

O'REILLY®
Technical Guide

Accelerating Data Pipeline Development

Deliver Data Projects Faster
Without Creating Tech Debt

Compliments of



coalesce®

Josh Hall



Deliver data projects 10x faster – without creating tech debt

Build a supercharged data foundation
and cut through the noise of data
development with Coalesce.

The screenshot displays the Coalesce web application interface. At the top, there's a navigation bar with the Coalesce logo, a search bar, and links for Build, Deploy, and Docs. Below this, the main interface is divided into several sections:

- Nodes:** A sidebar on the left lists various data nodes under 'Source' and 'Stage' categories. The 'Source' category includes COUNTRY, CUSTOMER_LOYALTY, FRANCHISE, LOCATION, MENU, ORDER_DETAIL, ORDER_HEADER, and TRUCK. The 'Stage' category includes STG_COUNTRY, STG_CUSTOMER_LOYALTY, STG_FRANCHISE, STG_LOCATION, and STG_MENU.
- Table View:** The central part of the interface shows a table for the 'DIM_CUSTOMER_LOYALTY' node. It has columns for Column Name, Transform, Data Type, Source, Nullable, and Description. The table lists various attributes like CUSTOMER_ID, FIRST_NAME, LAST_NAME, CITY, COUNTRY, POSTAL_CODE, PREFERRED_LANGUAGE, GENDER, FAVORITE_BRAND, MARITAL_STATUS, CHILDREN_COUNT, SIGN_UP_DATE, BIRTHDAY_DATE, E_MAIL, PHONE_NUMBER, SYSTEM_VERSION, SYSTEM_CURRENT_FLAG, SYSTEM_START_DATE, SYSTEM_END_DATE, SYSTEM_CREATE_DATE, and SYSTEM_UPDATE_DATE.
- Config Panel:** On the right, there's a configuration panel for the selected node. It includes sections for 'Node Proper', 'Options', 'Create As', 'Multi Source', 'Business Key', and 'Channel Tracking'. The 'Create As' section shows a dropdown menu set to 'Table'. The 'Multi Source' section has a toggle switch. The 'Business Key' section shows a list of 12 items with search and selection buttons. The 'Channel Tracking' section shows a list of 13 items with search and selection buttons.

Discover the future of data
transformations at Coalesce.io

Accelerating Data Pipeline Development

*Deliver Data Projects Faster
Without Creating Tech Debt*

Josh Hall

O'REILLY®

Accelerating Data Pipeline Development

by Josh Hall

Copyright © 2025 O'Reilly Media, Inc. All rights reserved.

Published by O'Reilly Media, Inc., 141 Stony Circle, Suite 195, Santa Rosa, CA 95401.

O'Reilly books may be purchased for educational, business, or sales promotional use. Online editions are also available for most titles (<http://oreilly.com>). For more information, contact our corporate/institutional sales department: 800-998-9938 or corporate@oreilly.com.

Acquisitions Editor: Aaron Black
Development Editor: Melissa Potter
Production Editor: Christopher Faucher
Copyeditor: Paula L. Fleming

Proofreader: Arthur Johnson
Interior Designer: David Futato
Cover Designer: Ellie Volckhausen
Illustrator: Kate Dullea

September 2025: First Edition

Revision History for the First Edition

2025-09-12: First Release

See <http://oreilly.com/catalog/errata.csp?isbn=9798341608740> for release details.

The O'Reilly logo is a registered trademark of O'Reilly Media, Inc. *Accelerating Data Pipeline Development*, the cover image, and related trade dress are trademarks of O'Reilly Media, Inc.

The views expressed in this work are those of the author and do not represent the publisher's views. While the publisher and the author have used good faith efforts to ensure that the information and instructions contained in this work are accurate, the publisher and the author disclaim all responsibility for errors or omissions, including without limitation responsibility for damages resulting from the use of or reliance on this work. Use of the information and instructions contained in this work is at your own risk. If any code samples or other technology this work contains or describes is subject to open source licenses or the intellectual property rights of others, it is your responsibility to ensure that your use thereof complies with such licenses and/or rights.

This work is part of a collaboration between O'Reilly and Coalesce. See our [statement of editorial independence](#).

979-8-341-60876-4

[LSI]

Table of Contents

Preface.....	v
1. Getting Started in Coalesce.....	1
The User Interface	1
Projects	5
Workspaces	9
Storage	12
Adding Users	15
Adding Data Sources	16
Building on Your Foundation	17
2. Coalesce Core Concepts.....	19
Column-Aware Architecture	19
Node Types	22
The Pipeline Development Approach	30
The Development Workflow	32
Knowledge Sync Complete	33
3. Building Data Pipelines in Coalesce.....	35
The Build Interface	35
Adding Data Sources	38
Adding Nodes to Your Pipeline	41
Data Transformations in Coalesce	55
Joins	58
Bulk Editing	61
Data Transformation in Process	64

4. Managing Data Pipelines in Coalesce.....	65
Managing Nodes Using Views	65
Subgraphs	70
Selector Queries	72
Column-Level Lineage	73
Version Control	76
Macros and Parameters	82
Testing	85
Jobs	86
Environments	87
The Makings of a Pro	90
5. Coalesce Security and Data Governance.....	91
Account Access	91
Role-Based Access Control	93
Coalesce SQL Execution	99
Documentation	100
Peace of Mind	103
6. Modeling Patterns and Use Cases in Coalesce.....	105
Migration: From Legacy to Modern	106
Modeling	109
Wrapping Up	112
7. Wrapping It All Up with Coalesce Catalog.....	113
Getting Set Up with Coalesce Catalog	114
Metadata Management and Data Lineage	116
Collaboration and Data Discovery	119
Using the AI Assistant for Data Discovery and Exploration	122
From Chaos to Clarity	123

Preface

If you work in data today, you know the landscape is changing fast. Cloud platforms evolve, teams shift, and expectations continue to rise. Yet, for many data professionals, the day-to-day still feels frustratingly manual, brittle, and opaque. You might spend hours tracking down the source of a single metric. Or rewriting the same SQL logic across projects. Or wondering if a change will break downstream dashboards. Too often, data work is reactive, patchy, and invisible—until something goes wrong.

This is your guide to doing data work differently. It introduces Coalesce, a data transformation platform built for the modern data stack. But more than that, it introduces a new way of working that is structured, scalable, and transparent.

You won't find marketing fluff or generic best practices here. Instead, you'll find hands-on guidance rooted in real-world workflows. You'll learn how to build modular pipelines, enforce standards without sacrificing flexibility, and surface data that users can trust. Whether you're a solo analytics engineer or you're leading a large team, this guide will help you build data systems that scale with confidence.

What This Guide Covers

Throughout the guide, we'll be using the analogy of building a house to guide the process of developing data products in Coalesce. You'll begin by laying the groundwork: setting up your first Coalesce project, connecting to your data platform, and configuring your workspace. The first chapter will establish the core components

of development in Coalesce, such as projects, workspaces, storage mappings, and source data management.

Once the foundation is in place, you'll explore what makes Coalesce different. In the second chapter, you'll learn about column-aware architecture and how it enables automated lineage, bulk operations, and metadata capture at scale. You'll understand why Coalesce favors reusable logic and standardized development over ad hoc SQL scripts. These ideas aren't just technical; they're foundational to building maintainable, high-trust pipelines.

Next, this guide walks you through full pipeline development. You'll learn how to:

- Add data sources and build transformation logic using reusable node types
- Apply both column-level and node-level transformations
- Manage joins, build transformation layers, and document your work as you go
- Extend functionality in your projects with Coalesce Marketplace

As your pipelines grow, Coalesce helps you keep them organized. You'll use features like the node grid, the column grid, and sub-graphs to visualize your work. You'll build workflows using version control, jobs, environments, and parameter-driven deployments. You'll see how column-level lineage lets you trace transformations with precision and how the Problem Scanner helps you catch issues before they break any of your pipelines.

From the beginning, security and governance are baked into the product. You'll learn how to configure role-based access control (RBAC), enforce multi-factor authentication (MFA), and manage users across your organization. More importantly, you'll see how security integrates with development, not as a barrier but as a structure that promotes safety and autonomy.

Then the guide shifts focus to advanced modeling techniques and architectural patterns. You'll explore use cases like:

- Migrating from legacy pipelines to modern cloud platforms
- Designing clean, scalable dimensional models for business intelligence (BI)

- Implementing a data vault for full historical traceability
- Operationalizing AI and ML pipelines with nodes

Each use case highlights how Coalesce balances structure with flexibility, letting you build systems that are powerful but still easy to manage and understand.

Finally, the guide closes with a discussion of Coalesce Catalog, the capstone that connects everything you've built. Catalog takes your lineage, documentation, ownership, and usage signals and makes them accessible across your company. It turns your data pipelines into something discoverable, navigable, and genuinely useful. Whether you're an engineer or a business stakeholder, Catalog gives you the tools to explore data confidently and get answers without relying on institutional knowledge.

Why Coalesce? Why Now?

The modern data stack has delivered huge advances: cloud native warehouses, orchestration tools, reverse ETL, and more. But while storage and compute have become easier and more scalable, transformation processes often remain messy and difficult to manage. Codebases can become cluttered and hard to navigate. Documentation is frequently disconnected from the pipeline, making it difficult to understand the assets you're working with. Debugging relies heavily on institutional knowledge or a deep familiarity with legacy logic. As tools and teams have evolved, collaboration has often been left behind, making data work feel isolated and reactive.

Coalesce was built to change this. It offers a unified experience where you can:

- Design transformations visually and write SQL when needed
- Define and reuse data patterns using nodes and templates
- Automatically track object- and column-level lineage and impact
- Use built-in testing, scheduling, and deployment workflows
- Manage access and document your work without extra overhead

Most importantly, it lets data teams work like software teams: with structure, visibility, version control, and automation. But unlike code-heavy frameworks, Coalesce doesn't require you to be a full-stack engineer to be effective. It supports fast iteration for analysts, best practices for engineers, and clarity for stakeholders—all in one platform.

You'll see these principles in action throughout the guide. Each chapter builds on the last, gradually expanding from tactical development to broader architectural thinking. By the end, you won't just know how to use Coalesce—you'll know how to use it to lead high-functioning data projects.

Who This Guide Is For

This guide is for anyone responsible for turning raw data into something reliable, scalable, and actionable, whether you're an analytics engineer focused on building faster, more maintainable pipelines, a data engineer aiming to standardize and automate workflows at scale, a BI developer or analyst trying to trace logic and lineage with clarity, or a platform or governance team working to establish sustainable data practices across an entire organization.

You don't need to be an expert in any one domain. In fact, Coalesce is often most powerful for cross-functional teams, where analysts, engineers, and product owners need to collaborate around a shared pipeline.

What matters is that you care about building data products right. You've probably felt the pain of unmaintainable logic, broken dashboards, or slow handoffs. You might be looking for a better way to scale your work or share knowledge across teams. This guide will help you do all of that.

Whether you're a team of one or you're leading an enterprise data platform, the goal is the same: clarity, consistency, scalability, and speed. And that's exactly what Coalesce is designed to deliver.

A Final Word Before You Dive In

This isn't about learning another tool for the sake of it. It's about changing the way you think about data work. Data pipelines aren't just technical assets. They're *processes*. They reflect your business logic, your priorities, and your assumptions. If they're messy, opaque, or brittle, so is your understanding of the business. But when they're clear, well documented, and modular? They become assets your entire organization can trust.

Coalesce helps you build pipelines like that, in a fast and sustainable manner. Not just accurate, but explainable; and not just for today, but in a way that scales with your team.

You bring the business logic, and Coalesce gives you the tools to express it clearly. Let's get started.

Getting Started in Coalesce

Just like the foundation that supports a house, Coalesce has foundational components that support your development of data products. This chapter will lay the groundwork for the rest of this guide, ensuring you have the tools you need to build end-to-end data pipelines. In it, you will learn about projects and workspaces, which are the core of the development experience in Coalesce. You will discover how to connect Coalesce to your data platform and make your data ready for development. You'll also see how you can work with the rest of your data team to build pipelines for all your data transformation needs.

Equipped with this knowledge, you'll be able to get started developing your data in Coalesce and have a foundational understanding for the rest of the topics discussed in this guide. Let's start by learning about the user interface.

The User Interface

Coalesce provides a graphical user interface (GUI) that enables you to develop your data while giving you the flexibility to write Structured Query Language (SQL). The interface is divided into different segments that provide support for different functions throughout the data development lifecycle. These segments include:

- The projects page
- The build interface
- The deploy interface
- The docs interface

You also have the ability to manage your Coalesce organization and users from the user menu interface using Organization and User settings. In this section, you'll explore each of the segments of the Coalesce interface.

The Projects Page

The *projects page* is the default landing page when logging into Coalesce. This page will display any of your organization's projects that you have access to. Projects in Coalesce give you the flexibility to organize your data initiatives for a specific purpose or team goal, as shown in [Figure 1-1](#). Don't worry if you can't make out all the details here; we'll dig into these screens in more detail shortly.

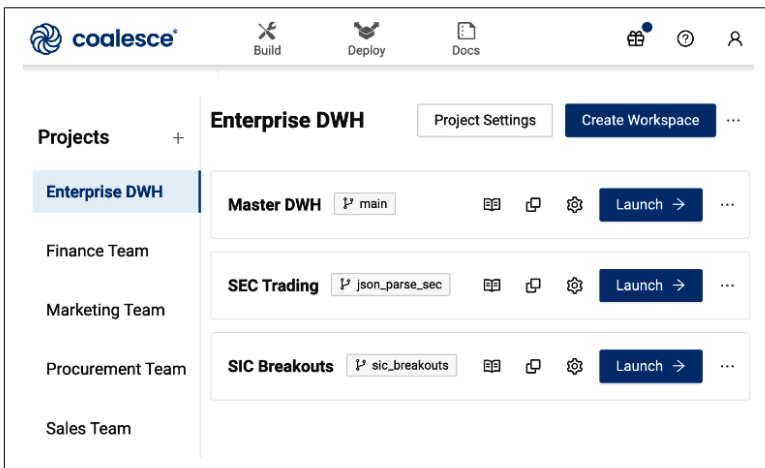


Figure 1-1. The projects page of Coalesce

The Build Interface

The *build interface* is where you will develop your data products and build your graphs and pipelines, as shown in [Figure 1-2](#). It can be accessed by launching any workspace from the projects page. Users can easily manage each aspect of a data pipeline from the build interface, including creating jobs and subgraphs.

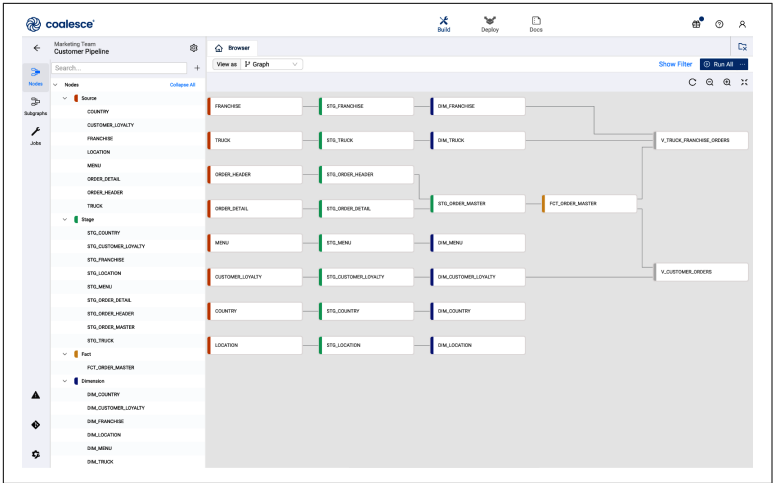


Figure 1-2. The build interface in Coalesce displaying various objects

The Deploy Interface

The *deploy interface*, shown in **Figure 1-3**, is where you can deploy your data projects from the desired state in your Git repository. You can also see a history of your pipeline's deployments and refreshes from the feed on the right side of the screen. This interface contains a dashboard for any environment you create, allowing you to see the current status of job runs as well as to schedule job refreshes.

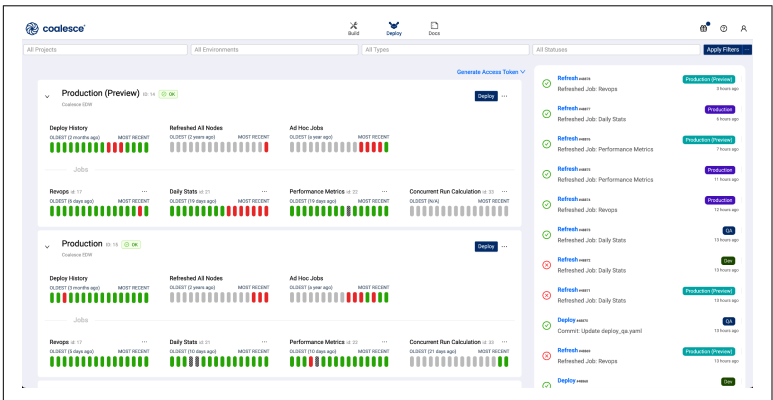


Figure 1-3. The deploy interface, where you can manage and monitor your data projects

The Docs Interface

The *docs interface* captures automatic, real-time documentation about each project and environment in your Coalesce organization. Here you can find information about each project created, such as the database and schema, column names and descriptions, and even data definition language (DDL) and data manipulation language (DML), as shown in [Figure 1-4](#).

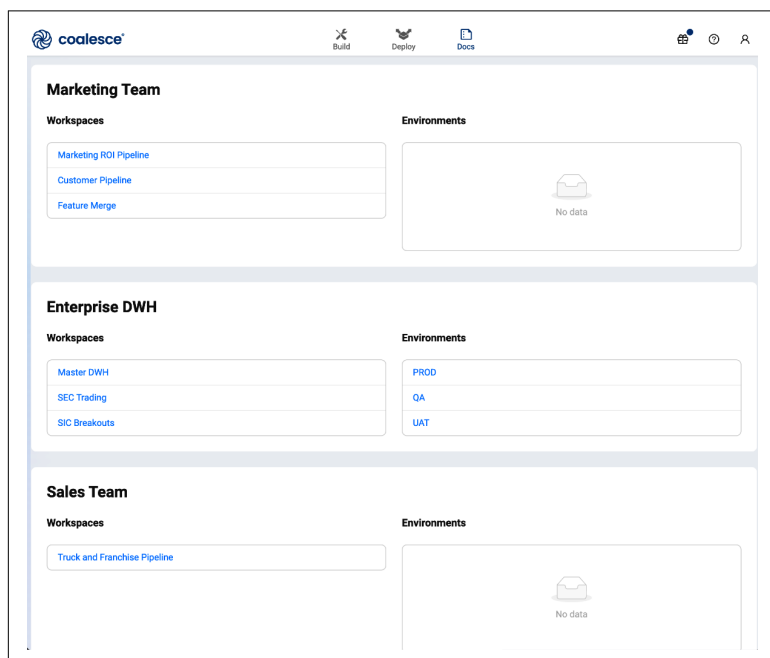


Figure 1-4. The docs interface displaying all of the workspaces and environments that have been documented

The User Menu

As a user within Coalesce, you will have access to the *user menu*. The user menu can be accessed via the user icon (the outline of a person) in the upper right-hand corner, as shown in [Figure 1-5](#). Within the user menu are the Organization and User settings of your organization.

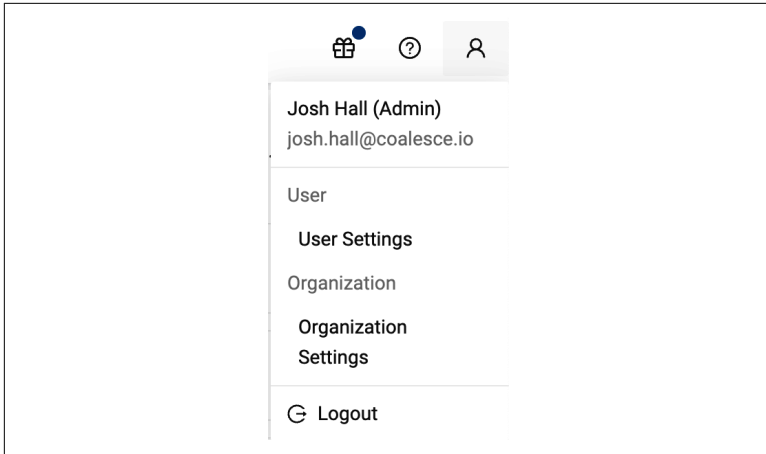


Figure 1-5. The user menu opened to display User and Organization settings

Organization settings provide management access to the following controls for your Coalesce organization:

- User management: Adding, removing, or modifying users
- Single sign-on
- Preferences, such as the Coalesce parser sample size

User settings provide you with the ability to manage your individual user experience with the following:

- Configuration of your Git settings
- Support information about your Coalesce account
- Password changes

With the ability to navigate the Coalesce interface, you can now dive into setting up your first Coalesce project for developing your data.

Projects

You learned about the projects page in the previous section, but now it's time to take a step further to learn more about projects. *Projects* give you the ability to organize your data development in a structured way, similar to how folders allow you to organize documents in Google Drive. In this section, you will learn how to

use projects, as well as how to set up a project for your own data development.

The Purpose of Projects

Projects provide multiple advantages for data teams developing their data. The first of these advantages is architectural. By utilizing projects, you can decide how your data development processes will be architected.

For some data teams, this means organizing projects according to the initiatives that the data team is working on. Others may want to implement a data mesh pattern and use projects to separate each domain of the business. Still others may want different teams within the organization to be managed through separate projects, as shown in [Figure 1-6](#). Regardless of the organization pattern, you must have at least one project in order to develop your data.

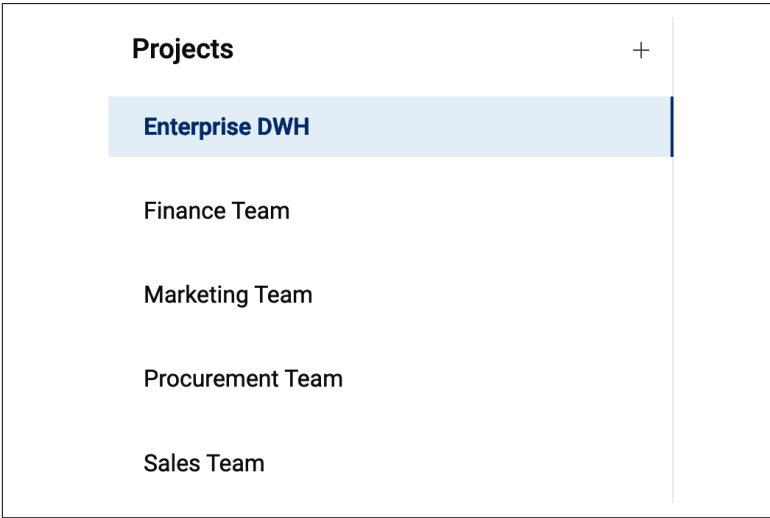


Figure 1-6. Projects in Coalesce managed by the team working on the project

Another advantage of projects is version control. Each project in Coalesce is integrated with your version control provider, such as GitHub. You will integrate a Git repository for each project that you create. This allows each project to be managed separately from the others while providing version control for the specific purpose of the project.

While it is possible to skip this integration, your development experience will be limited to a singular workspace. Coalesce does not recommend developing this way.

The last advantage we'll discuss is role-based access control, or RBAC. RBAC allows your Coalesce administrators to determine which users should have access to which projects, as well as which level of access each user should have. This provides granular control for all of your data development initiatives.

How to Set Up a Project

Now that you know when to use a project, let's talk about how to set one up. If you have never configured a Git account in Coalesce before, you will need to do this first. To configure a Git account in Coalesce, navigate to the User settings. Within the version control section, select Add New Account. You can follow the instructions provided by the Git setup modal to enter the information necessary to configure your Git account, as shown in **Figure 1-7**.

Add Version Control Account

* Account Nickname:

A friendly nickname for this account, e.g. 'GitHub'.

Credentials

* Username:

A user-friendly name for the version control account, typically your account username or name.

* Token:

A personal access token or application password.

Author Details

Identifies you as the author of your commits.

* Name:

* Email:

Cancel

Add Account

Figure 1-7. The Git setup modal in Coalesce

From the projects page, click the plus (+) button next to the Projects header. This will open the project configuration workflow. Here you'll give your project a meaningful name, select the data platform you want to connect Coalesce to, and provide any descriptive information, as shown in [Figure 1-8](#).

Step 1 of 3

Create a Project

Project Details

* Name

Enterprise Design Team

* Platform

Snowflake

Description

This project is dedicated for the enterprise design team to focus on menu development and understanding user behavior resulting from menu design

Previous Next

Figure 1-8. The project setup workflow, where you select your data platform and supply a name and description

Next, you'll need to provide the Git repository URL from your version control system. You can skip this step to create a project without version control, but I don't recommend this. If you choose to skip this step, you won't be able to save or version any of your work. Once you supply the Git repository URL, you can select a Git account configuration from Coalesce. Once you have selected your Git account, your setup is complete, and you can complete the project setup workflow.

Workspaces

With your new knowledge of projects, you can now move on to setting up a workspace for your data development projects. A *workspace* is a sandbox environment where you can complete the development of your data. Each workspace has its own graph, storage locations, macros, node types, data platform connection configuration, and Git branch—you'll learn about all of these in subsequent sections. You can create multiple workspaces to work on different tasks and merge them into your codebase. To begin building your data pipelines, you will need to connect your workspace to your data platform. Let's go over how to do that now.

Connecting a Workspace to the Data Platform

Within any project, select the Create Workspace button in the upper right corner of the project to open the workspace creation workflow. You will be asked for a workspace name and description. Next, you will select the branch and commit you want to create your new workspace from. For example, you may want to create a workspace from your main branch in order to design a new forecasting pipeline in your data warehouse. Once you have selected the branch and commit, you will supply the name of the new branch you wish to create and finish the creation of the workspace.

With a workspace created, you can now connect it to your data platform. By clicking on the gear cog icon next to the workspace Launch button, you can provide the information about your data platform connection. In **Figure 1-9**, the workspace is connected to Snowflake, but Coalesce supports a number of different data platforms.

Coalesce also supports multiple authentication types. The two most common are OAuth and Username and Password. Once you have supplied Coalesce with your account information and connection criteria, you can test the connection to your data platform. Once successful, you are ready to map your workspace to the data that exists in your data platform.

Workspace Settings - Customer Pipeline

×

Settings

User Credentials

Storage Mappings

Parameters

OAuth Settings

Connection

Snowflake Account

https://fka56740.snowflakecomputing.com

Accounts are configured under Settings

* Authentication Type

Username and Password (Cloud)

▼

* Username

JHALL

* Password

Edit

Role

SYSADMIN

Leave empty to use the Snowflake default.

Warehouse

COMPUTE_WH

Leave empty to use the Snowflake default.

Test Connection

Delete Connection

Cancel

Save

Figure 1-9. Connecting your workspace to your data platform, in this case Snowflake

Once you have connected your workspace to your data platform, you can select the blue Launch button to enter the build interface. Before you can begin your data development in the build interface, however, there are a few more configuration items to complete, which you will find in the Build and Workspace settings. We'll quickly explore each of these.

Build Settings

The Build settings page can be accessed by clicking the gear cog icon in the lower left corner of the build interface. Within the *Build settings*, you can manage everything related to development within your workspace. This includes:

- Storage locations and storage mappings
- Development workspaces
- Environments
- Macros
- Node types
- Packages

Each of these items has its own settings, which can be configured by selecting the item.

Workspace Settings

Your *Workspace settings* allow you to manage the connection to your data platform. These settings can be accessed in one of two ways: you can either select the gear cog icon next to the workspace name or select the same icon next to the workspace in the Workspace selection from the Build Settings, as shown in **Figure 1-10**.

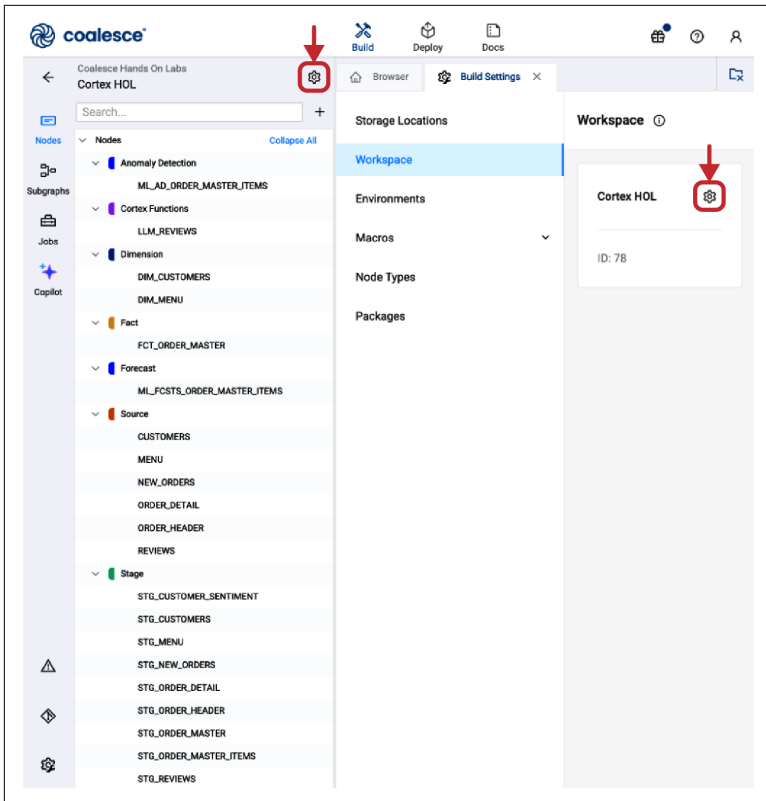


Figure 1-10. Where to open the Workspace settings inside of a workspace

Now that you know where all of the settings are in your workspace, you can complete your workspace configuration to begin building your data pipelines.

Storage

So far, you have created a project and workspace and connected the workspace to your data platform. But now you need to tell your workspace what data in your platform you want to develop with. This is where storage locations and storage mappings come into play.

Storage Locations

Storage locations are a logical representation of a database and schema in your data platform. You can think of them as the glossary that points to the chapters in a book. Storage locations themselves don't contain any data or perform any action on their own. Instead, they act as logical containers for the databases and schemas you want to use in your data pipeline development.

For example, you may have a database and two schemas within that database that contain source data for a pipeline you wish to develop. You want the results from your data pipeline to be output to a different database and schema. In this case, you could use three storage locations, two for the source data and one for the output or target destination, as shown in [Figure 1-11](#).

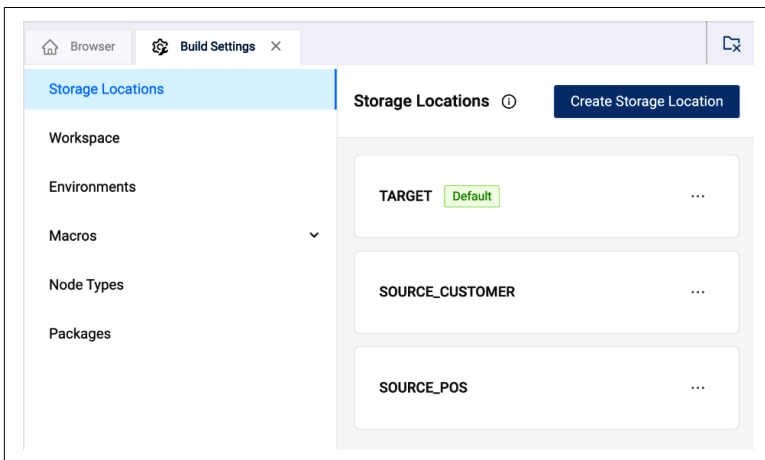


Figure 1-11. Storage locations created to be mapped to physical destinations in your data platform

Storage locations can be used for other use cases as well. For example, if your organization uses a medallion architecture, you can have a storage location for each level of your architecture—e.g., bronze, silver, and gold. Or if you leverage a staging layer in your data pipelines, you could have a staging storage location where all staging tables are created in the same location.

You will notice that there is always a *default* storage location. This is the location that, unless configured otherwise, all tables in your data pipeline will be created within by default.

It's important to note that *storage locations cannot be renamed once they are created*.

Once your storage locations are created, you can map each one to a physical destination in your data platform using storage mappings.

Storage Mappings

As you just learned, storage locations are just logical containers that can point to physical locations in your data platform. *Storage mappings* are what tie your storage locations to a database and schema in your data platform. To configure storage mappings for your workspace, you will need to access your Workspace settings and select Storage Mappings.

Each of the storage locations you have created will show up as an item to be mapped to your data platform. For each storage location, you can select the database and schema that you want each to point to, as shown in **Figure 1-12**. Ensure that your default storage location is mapped to a database and schema where you expect your tables to be output to.

		Database	Schema
TARGET		JHALL_DEV	PUBLIC
SOURCE_CUSTOMER		FROSTBYTE_TASTY_BYTES	RAW_CUSTOMER
SOURCE_POS		FROSTBYTE_TASTY_BYTES	RAW_POS

Figure 1-12. Mapping physical databases and schemas to the storage locations previously created

Storage mappings can be updated at any time from the Workspace settings. You can also add more storage locations and map them in the same way described here. For this guide, we will be using a fictional food truck company dataset in our examples. Our data

sources will be the customer and point-of-sale (POS) data, and our target storage location will be a development database and schema. We can deploy this to a production environment later on.

Adding Users

You now have a fully configured and ready-for-development project and workspace. With your data ready for development, you can add your team to your Coalesce account to collaborate alongside you. To add users, you will need to access the Organization settings from the user menu. As a Coalesce administrator, you will be able to add users to your Coalesce account and assign them appropriate roles.

Table 1-1 provides a guide to user roles in Coalesce and their associated permissions.

Table 1-1. Coalesce user roles and associated permissions

Role	Permissions summary	Recommended for ...
Organization administrator	The creator of the Coalesce app is automatically assigned the organization administrator role. Administrators have full access to all functionality in Coalesce. Only administrators can add other users, including other organization administrators.	Full administrative control
Organization contributor	Organization contributors have access to read documentation and to create API tokens, user settings, and Git account information. They can set up a project, configure Git, add members to projects, and oversee work. However, they have access only to the projects they create themselves. If there are multiple organization contributors, they will need to share access with each other. Contributors cannot add new users to the organization.	Managers who decide how each person will contribute to a project
Organization member	This is the default role. Organization members can edit Git account information, create API tokens, and read documentation.	Default role

Now that your users are added to your workspace with the proper permissions, your team can begin building data products by adding data sources.

Adding Data Sources

With your team ready to collaborate in Coalesce and your workspace ready to develop your data, you can begin adding data sources to your graph. You can do this by launching your workspace and selecting the plus (+) button in the upper left corner of the build interface. This will open the add data sources modal, which will display all of the storage locations you mapped your data to and the objects available in those locations, as shown in **Figure 1-13**.

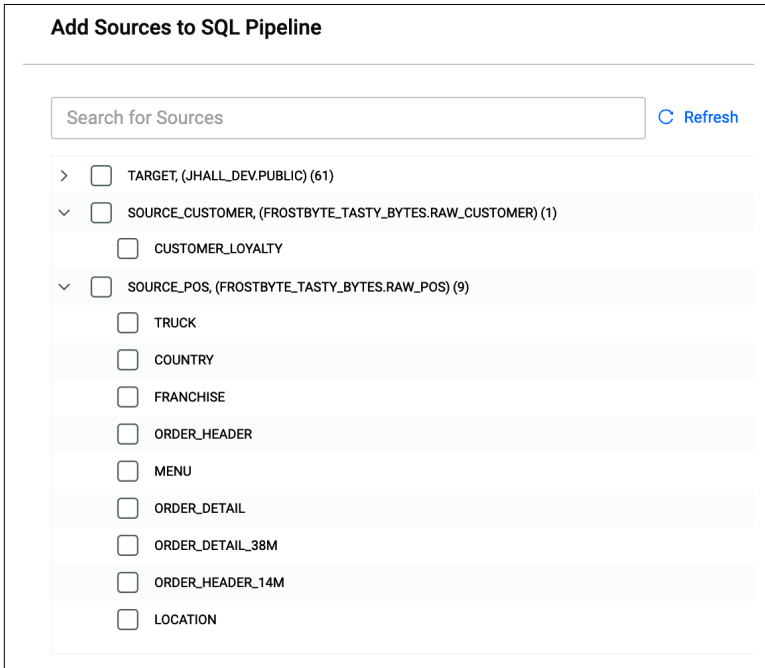


Figure 1-13. Add data sources modal showing all of the objects available from each of the configured storage locations

You can select as many or as few data objects as you want to add to your data pipeline. Once you have selected the data sources you want to work with, you can add them to your graph. Coalesce will add each object to your graph as an independent object. In the next chapter of this guide, we will discuss these objects in detail.

Building on Your Foundation

This chapter laid the foundation for all the components you will be using in Coalesce. In it, you learned how to navigate the user interface and how each segment of the interface is used. You learned about the purpose of projects and how to set one up. You also learned how to create a workspace and connect it to your data platform. At the end of the chapter, you saw how to bring data sources into your graph. At this point, you are ready to begin developing your data. As you go through this book, don't hesitate to revisit this chapter to ensure your foundation remains strong as you build in Coalesce.

However, before you begin building a data pipeline, it's important to understand the core concepts of how to develop your data in Coalesce. In the next chapter, we'll explore everything you need to know to begin your journey building data pipelines and applying all of the knowledge from this chapter to the concepts of pipeline development in Coalesce.

Coalesce Core Concepts

In the previous chapter, we laid the foundation for your understanding of Coalesce components, such as projects, workspaces, and storage locations. In this chapter, we will set the building blocks on top of your foundation by providing you with the core concepts that empower Coalesce data development. You will learn why column-aware architecture is important within data transformations, what nodes are and how they harness column-aware architecture, the data pipeline approach, and managing data development.

At the end of the chapter, you'll walk away with the knowledge that will provide you with the framework you need to understand and develop your data on the platform. And there is no better place to start than understanding column-aware architecture.

Column-Aware Architecture

Data processing workloads in modern data platforms don't just operate on the scale of thousands of database tables; they run on hundreds of thousands—even millions—of columns. Coalesce is built from the ground up to automate data transformations while supporting both a code-first approach and an intuitive visual interface. This all starts with a column-aware architecture.

But what is column-aware architecture? *Column-aware architecture*, or being column-aware, is an approach to managing data transformation with an understanding of columns and how they are connected. With this understanding, Coalesce provides automated

column-level lineage while enabling the creation and maintenance of database objects at scale. This column-aware architecture captures metadata for each object created, allowing you to leverage incredible development speed and agility in your data transformation workloads.

Thanks to its column-aware architecture, Coalesce shifts the paradigm of traditional data transformation processes to an automated, reusable, and scalable approach. We'll see how this translates to the foundational building blocks of Coalesce (nodes), but first, let's dig deeper into the benefits of building a data transformation platform using column-aware architecture.

Data Patterns

Column-aware architecture is the key to standardizing how transformations are applied, how tables are structured, and how columns are logically connected. This standardization can be referred to as a *pattern*, which is a reusable step in a data transformation process that represents a logical transformation. This can include:

- Incremental and processing logic
- Materialization logic
- Deployment logic

Data patterns are essential for the creation, management, and accessibility of data, especially at enterprise scale. Coalesce provides a platform to rapidly implement data patterns by *leveraging metadata at the column level*. With column-aware metadata, you can build a single reusable data pattern that can be applied across any of your data transformations.

Take, for example, a simple *type 2 slowly changing dimension* (SCD)—an industry standard for tracking historical data, such as a customer's current address, their address six years ago, and every change in between. The complex SQL to deliver this functionality could require writing hundreds of lines of code. However, because Coalesce is column-aware, this logic can be defined once and then reused as many times as needed, without having to write the code each time.

In Coalesce, column-aware data patterns streamline complex, time-consuming manual coding tasks. This speeds up data processing and lets you focus on your broader strategy while applying reusable logic with confidence.

Impact Analysis and Lineage

Having data patterns is powerful, but if those patterns aren't coupled with the ability to understand and manage your entire pipeline in one place, you may end up spending all of the time you saved building patterns working to understand how changes impact your pipeline. This is a second powerful benefit of using column-aware architecture: it enables complete impact analysis and lineage at the column level.

Imagine you've just deployed a pipeline that powers critical stakeholder dashboards. It's Monday morning, and you wake up to an email from your SaaS ETL provider: they've changed the schemas of several tables in your pipeline. Panic sets in. Which columns in my pipeline are impacted? Are the dashboards already broken? How many transformations are affected?

Thanks to Coalesce's column-level lineage, you can quickly see how changes to your data affect everything downstream, as seen in [Figure 2-1](#). You'll know exactly what's been impacted and how each object and column affects your pipeline at any moment. There's no guesswork or managing dozens of SQL tabs, and you can fix issues before they become problems.

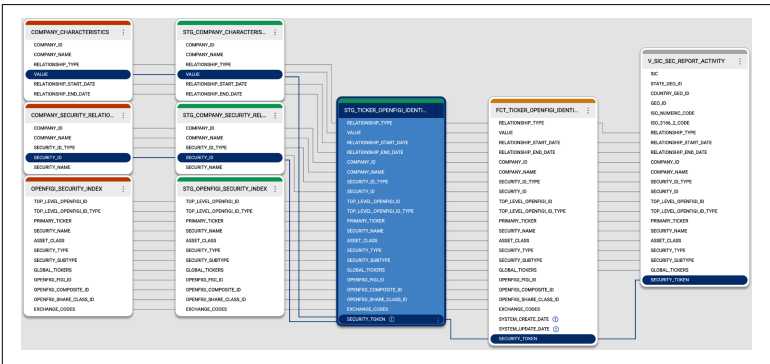


Figure 2-1. Column-level lineage showing the impact of every column throughout a pipeline

Additionally, you can perform bulk operations directly at column granularity and across your data platform, making source data changes or business logic changes easy and straightforward to implement. For a data consumer, Coalesce provides clear documentation at the column level, showing where the data came from and how it was calculated. This transparency helps build trust across the organization.

Scale and Governance

Column-level lineage drives efficiency and accuracy in data management—even across hundreds of thousands of columns, multiple teams, environments, and the entire business.

Column-aware *state management* minimizes the risk of data loss and errors by enabling in-place, column-level modifications instead of requiring full table re-creations. It also offers detailed column-level visibility into changes over time, a critical aspect of effective Data-Ops and governance practices.

For instance, in a deployment with thousands of columns, Coalesce eliminates the need to build complex architectures to handle changes. Instead, you can manage in-place edits out of the box. This approach reduces costs, enhances deployment visibility and planning, and fosters trust among data consumers throughout the organization.

Coalesce also offers column propagation, allowing you to add or remove columns across the entire pipeline with ease through the user interface. This functionality streamlines column management across your project.

Now that you have an understanding of column-aware architecture in Coalesce, let's dive into the foundation of data development: node types.

Node Types

Node types are the core components used to build data pipelines in Coalesce. A *node type* is a user-defined template of an object that exists in your data platform. Anytime we reference a node type in this guide, we are talking about the template that has been defined. You will also see references to a *node* or *nodes*, which are the visual representations of each node type in your graph, as

seen in [Figure 2-2](#). Each node type serves as a building block for constructing data pipelines and leverages data patterns to automate your data transformations.

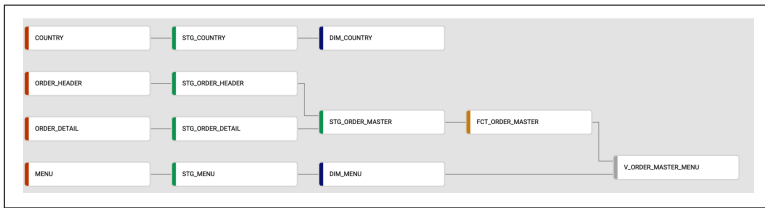


Figure 2-2. Nodes in the build interface representing a data pipeline

Node types enable pattern-based development by allowing you to define a transformation once and repeatedly apply it, streamlining automation and ensuring consistency across your pipeline. You can begin to see how leveraging nodes can accelerate development while ensuring quality and consistency across your pipelines, as everyone works from a shared pattern or standard. To build effectively with nodes, it's essential to understand how node types are created and function. Let's take a closer look at what makes up a node.

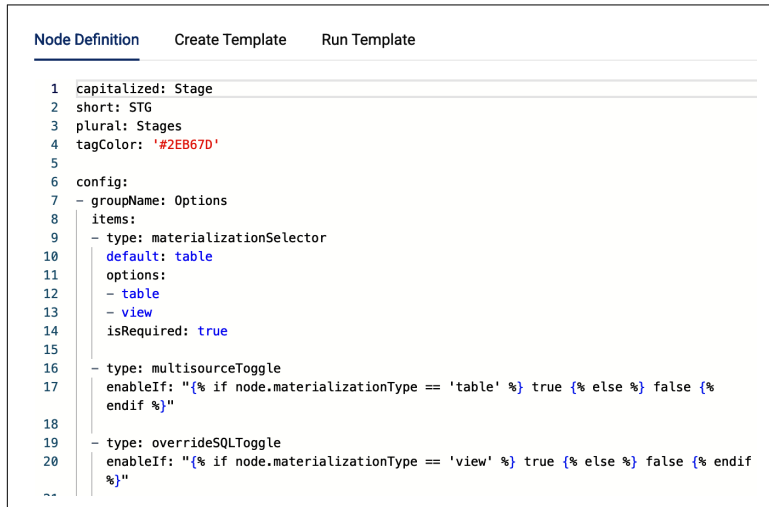
Node Type Architecture

Node types consist of three components: a node definition, a create template, and a run template. Each of these components performs a specific role in the representation and execution of a node type. Let's start with the node definition.

Node definition

The *node definition* defines the attributes available within a node type, such as naming conventions and node type colors. It also specifies the UI elements used to configure each individual instance of the node type. For example, if you want to build a Stage node type for creating a consistent staging layer in your pipeline, the node definition would define the naming convention for the node each time it's used. However, each instance of the Stage node can be configured differently, based on the options provided through the UI elements in the node's configuration. For example, one Stage node type could be materialized as a view and another could be materialized as a table.

The node definition is defined using YAML, as seen in [Figure 2-3](#).



```
1 capitalized: Stage
2 short: STG
3 plural: Stages
4 tagColor: '#2EB67D'
5
6 config:
7   - groupName: Options
8     items:
9       - type: materializationSelector
10         default: table
11         options:
12           - table
13           - view
14         isRequired: true
15
16       - type: multisourceToggle
17         enableIf: "{% if node.materializationType == 'table' %} true {% else %} false {%
18           endif %}"
19
20       - type: overrideSQLToggle
21         enableIf: "{% if node.materializationType == 'view' %} true {% else %} false {% endif
22           %}"
23
24 ..
```

Figure 2-3. YAML used in the node definition to define the attributes of the node type

The [Coalesce documentation](#) contains information on the configuration options available to include in your node definition YAML file, as seen in [Figure 2-4](#).



```
Node Definition Example

capitalized: 'My Node Name'           # Common Name (string, required)
short: 'MNN'                         # Node prefix. A '_' delimiter will be added auto
plural: 'My Node Names'              # plural name of common name (string, required)
tagColor: '#FF5A5F'                  # Node color. CSS colors or hex colors (string, r
deploymentStrategy: default          # Deployment strategy - see separate article for details.
isColumnResyncEnabled: true          # Allow non-source columns to be resynced.

config:                             # Array of the following config items
  - groupName: 'Core UI Elements'    # Name of config group (string, required)
    description: 'Core UI Elements'  # Description of group. Displayed in the GUI (str
    enableIf: 'true'                 # If true, display this config group, else hide t
                                     # Jinja expression resulting in a quoted boolean
                                     # (boolean string, 'true' | 'false', defaults 'tr

  items:                             # This will be followed by all your config items/

                                     ## Core UI Elements ##
  - type: businessKeyColumns          # Selector for business key columns
    displayName: 'Business Key'       # displayName for type businessKeyColumns (cannot
    attributeName: 'isBusinessKey'   # attributeName for type businessKeyColumns (cannot
    isRequired: false                # Require input from user for this config element
    enableIf: 'true'                 # If true, display this element, else hide this e
                                     # Jinja expression resulting in a quoted boolean
                                     # (boolean string, 'true' | 'false', defaults 'tr
```

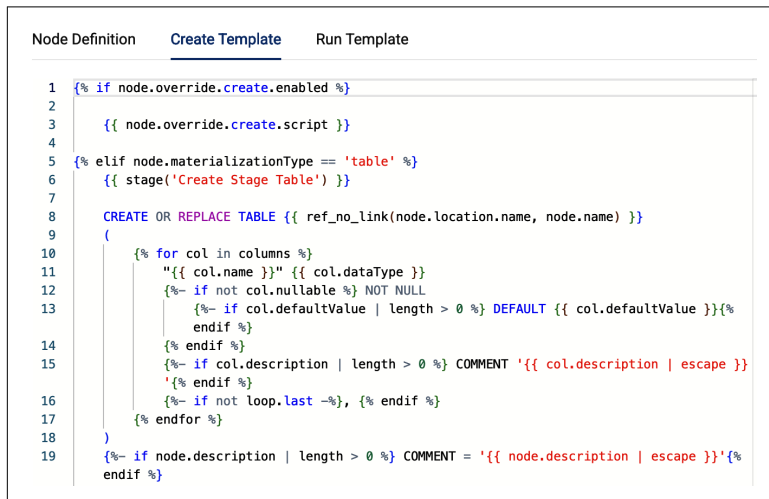
Figure 2-4. Coalesce documentation representing the various configuration options for the node definition of a node type

Once you have your node defined, you can configure the create and run templates.

Create template

The *create template* for a node type defines the logic used to automate the creation of the object. Typically, this involves DDL, which leverages attributes from the node definition—such as the selected materialization type—to execute a CREATE statement in SQL.

You can define create templates using a combination of SQL and Jinja, as shown in [Figure 2-5](#).



```
1 {% if node.override.create.enabled %}
2
3     {{ node.override.create.script }}
4
5 {% elif node.materializationType == 'table' %}
6     {{ stage('Create Stage Table') }}
7
8     CREATE OR REPLACE TABLE {{ ref_no_link(node.location.name, node.name) }}
9     (
10         {% for col in columns %}
11             "{{ col.name }}" {{ col.dataType }}
12             {%~ if not col.nullable %} NOT NULL
13             {%~ if col.defaultValue | length > 0 %} DEFAULT {{ col.defaultValue }}{%
14                 endif %}
15             {%~ if col.description | length > 0 %} COMMENT '{{ col.description | escape }}'
16                 {%~ if not loop.last -%}, {% endif %}
17             {% endfor %}
18         )
19     {%~ if node.description | length > 0 %} COMMENT = '{{ node.description | escape }}'{%
20         endif %}
```

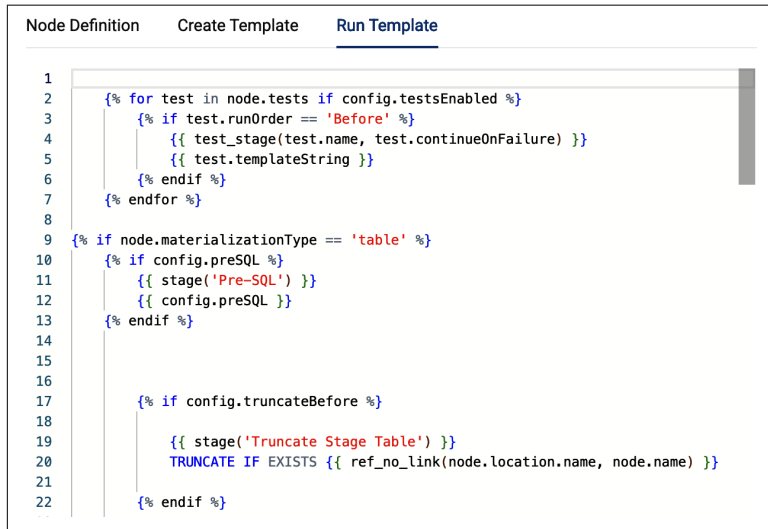
Figure 2-5. The create template representing the SQL and Jinja used to automate your DDL

After you define the logic, Coalesce runs the DDL for each object instance in your pipeline. Executing the create template creates the object in your cloud provider automatically. It is important to note that the create template is *only* creating the object; it is not loading data into the object. To insert data into the object, you need the run template!

Run template

The final component of a node type is the *run template*. This is where you define the logic to automate the execution of your transformations. Typically, it involves DML, which leverages Coalesce's column-aware architecture to execute SQL and populate objects with data.

Like create templates, run templates are defined using a combination of SQL and Jinja, as shown in [Figure 2-6](#).



```
Node Definition  Create Template  Run Template

1
2  {% for test in node.tests if config.testsEnabled %}
3    {% if test.runOrder == 'Before' %}
4      {{ test_stage(test.name, test.continueOnFailure) }}
5      {{ test.templateString }}
6    {% endif %}
7  {% endfor %}
8
9  {% if node.materializationType == 'table' %}
10    {% if config.preSQL %}
11      {{ stage('Pre-SQL') }}
12      {{ config.preSQL }}
13    {% endif %}
14
15
16
17    {% if config.truncateBefore %}
18
19      {{ stage('Truncate Stage Table') }}
20      TRUNCATE IF EXISTS {{ ref_no_link(node.location.name, node.name) }}
21
22    {% endif %}
```

Figure 2-6. The run template representing the SQL and Jinja used to automate your DML

With the logic defined, Coalesce will automatically run the corresponding DML for each node type. This will populate the objects in your data platform with data based on the logic and upstream data coming from your data pipeline. It's important to note that some node types may not need a run template. For example, a view node type is just storing a query to be run in the future; a view is never actually populated with data.

Now that you understand the components of a node type and how you can create reusable patterns, let's explore the advantages of developing with nodes.

Importance of Nodes

So far, we've seen how column-aware architecture drives pattern-based development, accelerating data transformations within Coalesce. Now we'll bring everything together to understand how these concepts translate into more efficient pipeline development.

Standardization

We've already discussed how node types establish a framework for consistent data development across your team. Now let's delve into why this standardization is essential for long-term, scalable pipeline development.

Standardization ensures that everyone operates from the same foundation, automatically and consistently. There's no need for new custom code, special permissions, or isolated environments. A node type is defined once, and it's ready to use repeatedly without additional setup.

Standardization enables your team to embed best practices directly into the nodes. Each time a node runs, you can trust it to perform exactly as intended. It also allows you to optimize your code, ensuring that each instance of a node type runs efficiently and minimizes compute resource usage within your data platform.

Standardization streamlines development by reducing repetitive object creation and manual coding. This allows your team to focus on data transformation and modeling rather than boilerplate setup. By minimizing development overhead, you can build and deploy pipelines more efficiently.

As discussed earlier, you can define a Dimension node type to support type 2 SCDs. Whenever you need to implement a type 2 SCD, you can add the Dimension node to your pipeline and configure it with just a few clicks, as seen in [Figure 2-7](#). In seconds, you can automate the execution of a type 2 SCD, eliminating the need to write code from scratch or manually adapt copied code to your specific tables.

Flexibility

Coalesce provides multiple built-in node types, including a Dimension node with preconfigured type 2 SCD options, to streamline development. While these predefined options simplify automation, data projects often require customization.

Coalesce supports both custom code and a GUI-driven interface (Figure 2-8), allowing you to define node behavior precisely while maintaining an efficient and user-friendly development experience.



Figure 2-8. Defining your own data patterns and writing custom SQL through a power user interface

Coalesce enables you to create user-defined nodes (UDNs), allowing you to configure nodes tailored to the specific needs of your data development. Additionally, Coalesce provides a variety of pre-built node types available from Coalesce Marketplace, designed to address a wide range of use cases. We'll learn more about Coalesce Marketplace in the next chapter.

Independent blocks of logic

You've likely encountered SQL transformations that span hundreds of lines, filled with multiple common table expressions (CTEs) and repetitive code that violates the DRY (don't repeat yourself) principle. Coalesce node types are designed to address this by breaking transformations into logical, reusable blocks.

Instead of bundling multiple CTEs and subqueries into one sprawling query, Coalesce allows you to handle each logical task in its own node. This means that each CTE in a large query can become a separate node in your pipeline.

You might be thinking that building pipelines this way will result in more objects than if you consolidate everything into a single

query—and you’re right! But this is only a problem when having more objects means more development time and management. Because Coalesce relies on standardized, reusable nodes, the building process is incredibly fast. And with column-aware architecture in place, managing a pipeline at any scale becomes simple. As we’ll explore in the next section, there are significant advantages to designing pipelines using nodes over relying on complex CTEs.

The Pipeline Development Approach

Complex data processing often involves breaking tasks into sequential steps, where the output of one step feeds into the input of the next. While CTEs can achieve this, Coalesce takes a modular, pipeline-based approach that offers significant advantages over traditional CTE development:

Transparency

One of the primary benefits of a pipeline approach is improved transparency. Instead of consolidating all logic into a single query, Coalesce uses nodes to represent individual logical blocks. This modularity makes it easy to navigate your pipeline and understand the function of each step. Additionally, you can run individual nodes independently to view their results in context, simplifying analysis and debugging.

Troubleshooting

The transparency of pipeline logic also simplifies troubleshooting. Debugging a monolithic SQL query with hundreds of lines and multiple CTEs can be tedious and error-prone. In Coalesce, each logical step is a standalone node, so identifying and resolving errors or data issues is a straightforward process. Thus, the modular design reduces the overhead required to troubleshoot and maintain data pipelines.

Testing

Coalesce facilitates out-of-the-box and SQL-based testing in each node in the pipeline, as shown in [Figure 2-9](#). For example, you can validate uniqueness, check for null values, or run other data quality checks at any step. By catching unsupported data or logical errors early in the pipeline, you can prevent issues from propagating to downstream processes, saving time and effort.

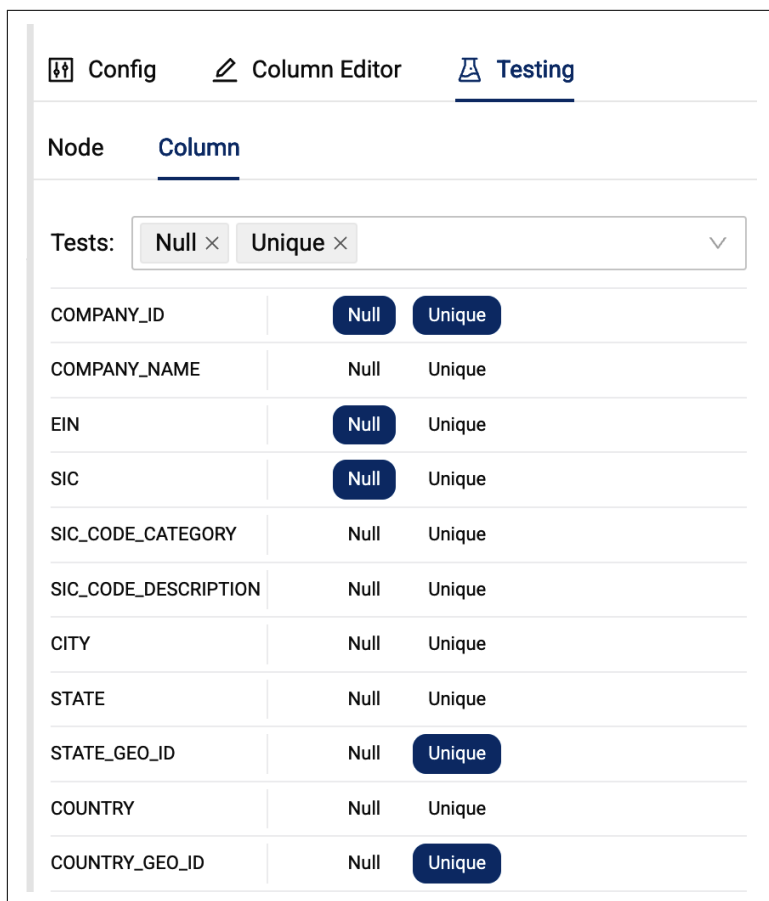


Figure 2-9. Out-of-the-box tests that can be applied easily to any column in the node

Developing pipelines in this way allows complete control over each element, while providing a streamlined approach to development. Breaking a data pipeline into nodes also promotes reusability of logic across the pipeline. For instance, if you create a node to deduplicate records in an orders table, downstream nodes can reference and reuse this logic without reprocessing. This eliminates redundant computation and ensures consistency across the pipeline. In contrast, using CTEs often requires duplicating logic across multiple queries, which is harder to manage and less efficient.

Developing with a pipeline approach offers substantial advantages in transparency, troubleshooting, testing, and reusability. While this method may result in more objects in your data platform, Coalesce’s column-aware architecture and node standardization make managing pipelines of any size seamless and efficient.

The Development Workflow

So far, we’ve explored the fundamentals of nodes and the advantages of building pipelines with them in Coalesce. But what does the actual development workflow look like? We’ll dive deeper into pipeline construction in the next chapter, but for now, there are two essential concepts to understand for developing your data in Coalesce—workspace development and environments:

Workspace development

Each workspace in Coalesce comes with a build interface designed for data development. When working within a workspace, each user will have their own set of credentials. Typically, development occurs in a sandbox or personal space within your data platform. This means that any objects you create are isolated to your storage location, allowing you to experiment and iterate without impacting shared environments.

Think of a workspace as your personal sandbox for developing and testing data pipelines. Once you’ve finalized your pipeline, you can deploy your work to an environment for broader use.

Environments

Coalesce provides environments to deploy data pipelines, allowing for a structured workflow from development to production. After finalizing changes in the workspace, commit them to a Git branch and deploy to the target environment.

Environments in Coalesce are tied to designated locations in your data platform, such as specific databases or schemas through storage locations and storage mappings. In keeping with best practices, Coalesce supports multiple environments, such as QA and PROD, allowing you to test and validate your work before deploying it to production. **Figure 2-10** illustrates an example of a typical environment setup.



Figure 2-10. The deploy interface in Coalesce, where you can deploy your pipelines to higher environments in your data platform

Knowledge Sync Complete

You are now equipped with the foundational knowledge needed to start developing your data in Coalesce. In this chapter, we covered column-aware architecture, how it supports pattern-based development through nodes, and why nodes serve as the optimal building blocks for building pipelines.

If you encounter challenges as you progress through this guide, feel free to revisit Chapters 1 and 2 for a refresher.

Next, we'll walk you through the process of building a data pipeline—from data source to insight-ready tables. See you in the next chapter!

Building Data Pipelines in Coalesce

So far, you've learned how to set up Coalesce and explored the core components and concepts of the platform, such as projects, workspaces, node types, storage locations and mappings, and column-aware architecture. In this chapter, you'll put this knowledge to work by learning how to build data pipelines in Coalesce.

You'll learn how to add data sources, create nodes, and write SQL to transform data directly within any node. You'll also master functionality like creating joins, bulk-editing columns, and writing filters. Plus, you'll discover how Coalesce Marketplace enhances your pipeline with powerful extensions.

Let's get started!

The Build Interface

In [Chapter 1](#), you got a brief introduction to the build interface, which is where you will be spending your time in this chapter. The *build interface* is where you will develop your data products and build node graphs and pipelines. You can access it by launching any workspace from the projects page, as seen in [Figure 3-1](#).

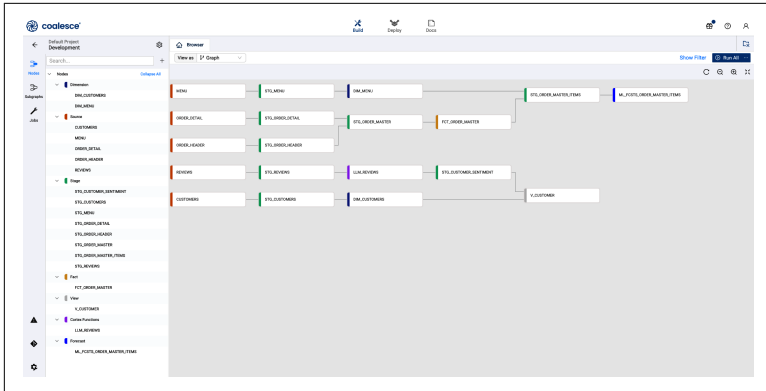


Figure 3-1. The build interface in Coalesce, displaying nodes organized in the form of a pipeline

The build interface contains all of the necessary components to build data pipelines. In the sidebar in the upper left side of the interface, you'll see navigation options for nodes, subgraphs, and jobs. You'll learn about subgraphs and jobs more in [Chapter 4](#), while the focus of this chapter will be on building pipelines with nodes.

In the lower left corner of the sidebar, you'll find navigation for your Problem Scanner, which will alert you to any issues or notifications you should be aware of within your data pipeline. You'll also be able to access your Git integration to easily version-control your work. Finally, you'll see the Build Settings cog, as seen in [Figure 3-2](#), which contains all of the settings for the workspace you are building in, such as the storage locations you learned about in [Chapter 1](#). Throughout the rest of this guide, we'll explore each of the items located in the Build settings, so don't worry if you don't know what each of these items is.

The build interface also contains the Workspace settings for the workspace you are working in. You learned all about the Workspace settings in [Chapter 1](#), when we discussed connecting to your data platform, but you can easily access those settings from the Workspace Settings cog in the upper left corner of the screen, as seen in [Figure 3-3](#). You can also find them in the Workspace line item in the Build settings.

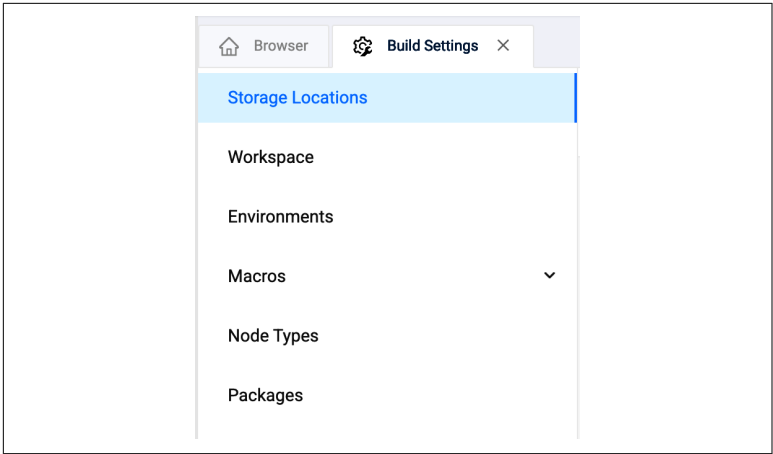


Figure 3-2. The Build settings available in each workspace in Coalesce

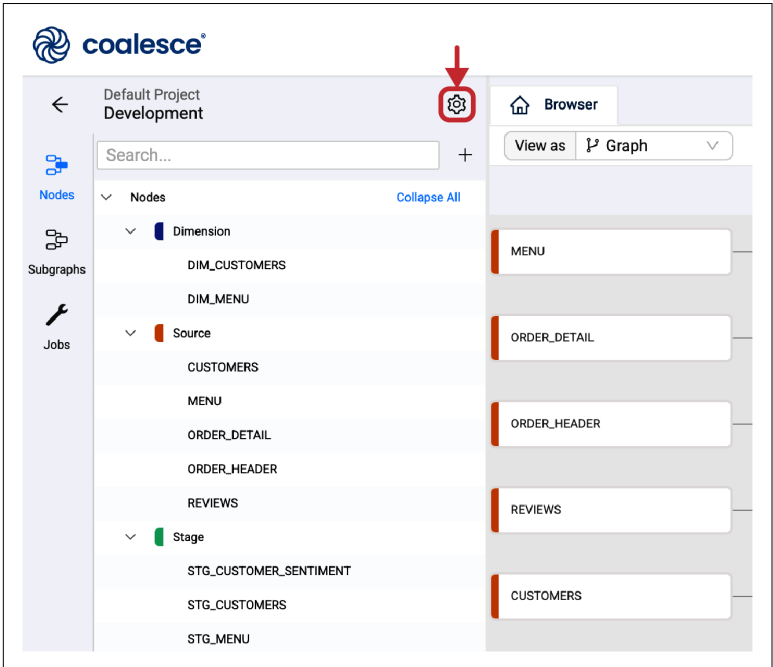


Figure 3-3. The Workspace Settings cog, which allows you to configure your Workspace settings

When it comes to building data pipelines, you'll do this work within the Browser, which is where your directed acyclic graph (DAG) and nodes in your pipeline are displayed, as shown earlier in [Figure 3-1](#).

Now that you're familiar with the essential elements of the build interface, it's time to start building data pipelines. First up: data sources.

Adding Data Sources

Data sources are a core component of any data pipeline—after all, you can't build a pipeline without a source. Adding data sources in Coalesce is straightforward. In the upper left corner of the build interface, select the plus (+) button and choose Add Sources. This will open the data sources modal, which will display all of the storage locations configured in the Workspace.

Each line item is a storage location that points to a database and schema. You can see in [Figures 3-4](#) and [3-5](#) how the storage mappings for each storage location correspond to the data sources available in the modal.

Storage Location Mapping

Override Mapping Values ☐

	Database	Schema
SAMPLE	SNOWFLAKE_SAMPLE_DATA	TPCH_SF1
WORK	PC_COALESCE_DB	PUBLIC
SALES	CORTEX_HOL	RAW_POS

Figure 3-4. Storage mappings in the Workspace settings, pointing your storage locations to the database and schema in your data platform where your data exists

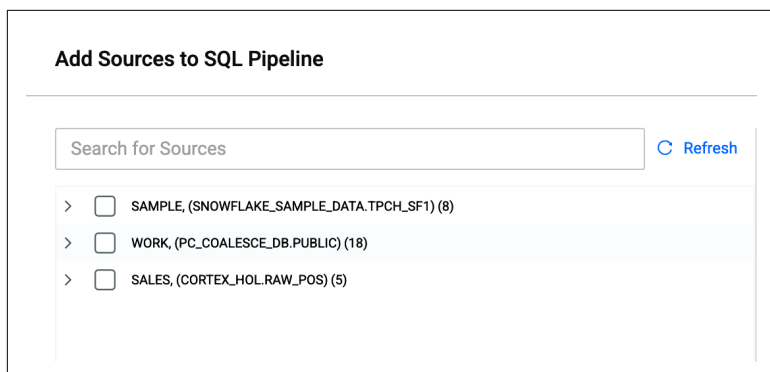


Figure 3-5. The same storage locations and mappings showing up in the data sources selector

By selecting any of the storage locations available, you can view all of the objects available to use within Coalesce. For example, in [Figure 3-6](#), you can see there are eight objects available that we could add into our pipeline.

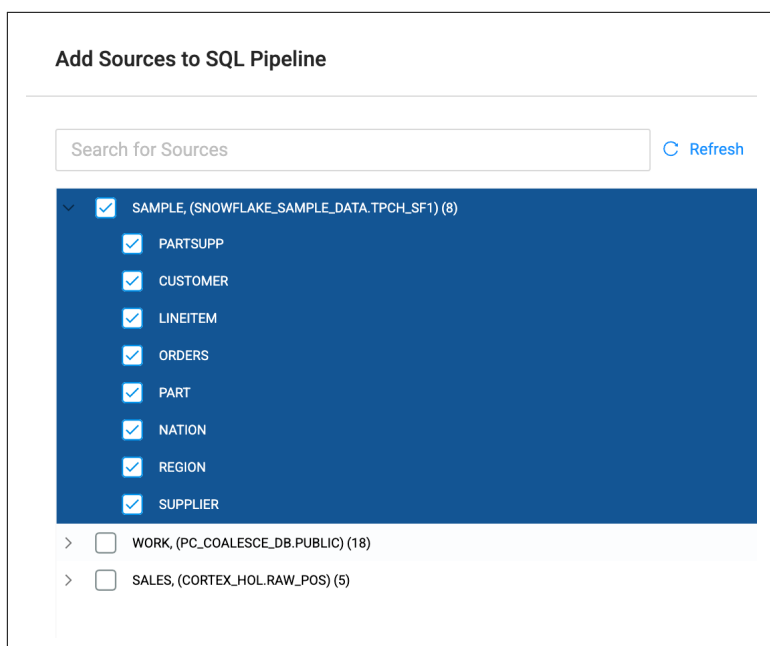


Figure 3-6. Selecting data sources from the data sources modal

You can either select the specific objects you need or select all of them by clicking the checkbox next to the storage location name. You can select as many objects as you want across any of the available storage locations. Once you have selected all of the objects you want to add to your pipeline as data sources, click the Add *X* Sources button in the lower right corner of the modal, where *X* represents the number of objects selected.

Coalesce will automatically add all of the objects you selected into the browser of the build interface, as seen in [Figure 3-7](#). Every kind of node has a unique color, and source nodes are dark orange. This will be consistent for any source node added to your Workspace.

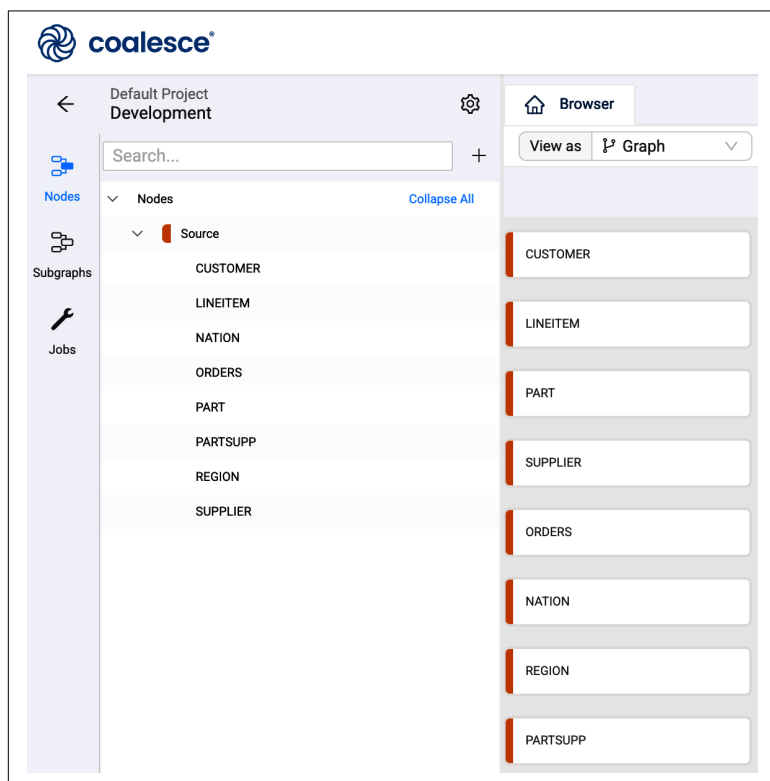


Figure 3-7. Data sources displayed in the Browser

With your data sources added to your Workspace, you can now begin transforming this data by using node types—whether out of the box, custom built, or, as we’ll see later in this chapter, from Coalesce Marketplace.

Adding Nodes to Your Pipeline

At the core of the developer experience in Coalesce is the ability to quickly add nodes to your data pipeline and start transforming data with ease. Nodes can be added directly from the platform or through Coalesce Marketplace. Because each node follows a standardized structure, the anatomy remains consistent across your pipeline, creating a predictable and unified experience for every developer. Coalesce includes several built-in node types to help you move faster right out of the gate:

- Stage
- Persistent Stage
- Dimension
- Fact
- View

You can right-click on any of the source nodes and hover over Add Node to view the node types, as shown in [Figure 3-8](#). Let’s explore each of the node types next.

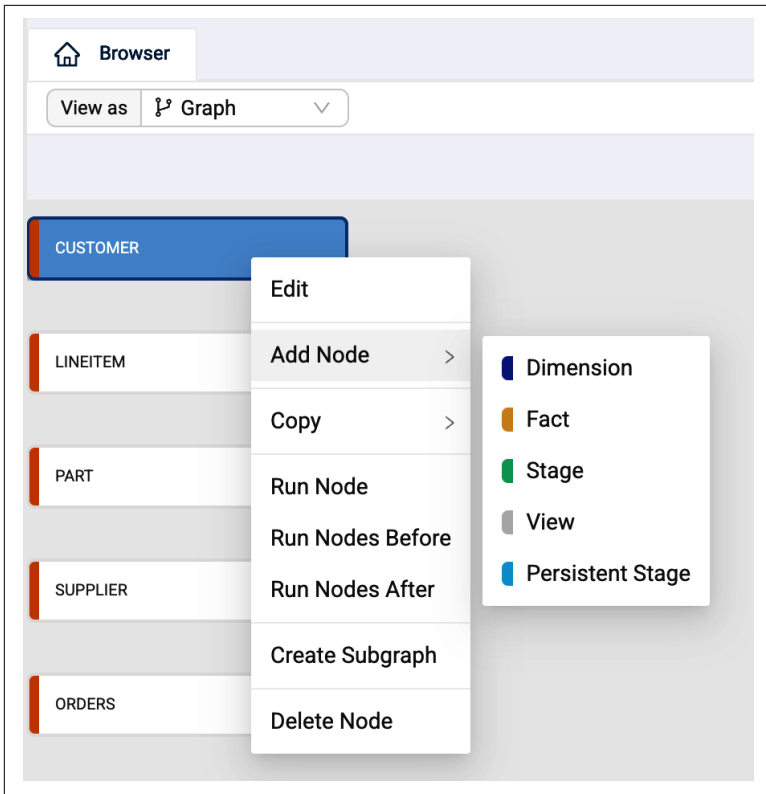


Figure 3-8. Node types available to add to a data pipeline

Stage

A *Stage node type* in Coalesce allows you to develop and deploy work in a table or view. It provides an intermediary working or staging layer to store, prepare, and transform raw data before downstream tables in your data pipeline use the data. Using a staging layer is a common data-engineering practice for preparing your data. Just as a cook needs to prepare raw ingredients such as carrots and onions by slicing and dicing them on a cutting board, you can think of Stage node types as a preparation space.

Stage node types are by default set to truncate the data in your table during each execution. This means that all data is deleted from the table and the fresh data coming from the upstream node is processed into the Stage node. If the truncation were deactivated, the node would append data to the underlying table.

Persistent Stage

Similar to a Stage, a *Persistent Stage* is an intermediary node type that provides data persistence across execution cycles. Unlike a Stage node type, a Persistent Stage contains a business key that lets you determine the unique identifier for your data source. This, in turn, allows you to persist or store historical data and load net new records into the object. This functionality is particularly beneficial when the objective is to retain historical data for prolonged periods.

Dimension

Coalesce supports type 1 and type 2 SCDs out of the box. Each *Dimension* requires a business key, or unique identifier, to be defined. You then have the option to define change-tracking columns, which will automatically configure the object as a type 2 SCD when columns are selected. This functionality is particularly beneficial for tracking changes to the dimensions (data that describes your facts) in your data.

For example, you may want to know anytime your customer's address changes when data is loaded from your sales data source. You can simply select the address column in your data source as a change-tracking column, and Coalesce will automatically generate best practice type 2 SCD Structured Query Language (SQL) within your data platform, saving you hours of development time.

Fact

Fact node types give you the ability to develop and deploy tables containing the measures or facts in your data platform. Each Fact node type contains a business key as well as functionality around varying operations for working with measures, e.g., revenue, cost of goods sold, profit, and so forth. By using Fact node types, you can easily distinguish your Fact and Dimension nodes in the Browser when looking at your DAG.

View

By default, the *View node type* is turned off when you are using Coalesce for the first time. You can enable it by going to the Build settings, selecting Node Types, and turning on the Enable toggle for the View node type. This node type allows you to create objects as a

view in your data platform, which means it is not storing any data. It also provides the flexibility to write custom SQL directly into the SQL editor.

While these five node types come out of the box, Coalesce is not limited to these. As we learned in [Chapter 2](#), you have the ability to create your own node types and, as we'll see next, use any of the nodes from Coalesce Marketplace.

Coalesce Marketplace

Coalesce Marketplace provides a wide array of packages that help bring extensibility to your data pipeline projects. These packages are composed of one or more node types that help you solve specific problems or provide functionality to accelerate data development, as seen in [Figure 3-9](#).

Package Marketplace

Featured Packages

Cortex

Leverage the power of Snowflake Cortex functions from within Coalesce.

[Find out more](#)

Dynamic Tables

Dynamic tables simplify data engineering in Snowflake by providing a reliable, cost-effective, and automated way to transform data.

[Find out more](#)

All Packages

Base Node Types

Coalesce base Node Types to construct facts and dimensions.

[Find out more](#)

Base Node Types Advanced Deploy

The Coalesce Advanced Deploy Nodes are Node Types that allows you to develop and deploy objects in Snowflake.

[Find out more](#)

Figure 3-9. Packages available on Coalesce Marketplace that can be added to your Workspace

Any package with the blue Certified checkmark means that Coalesce has certified the package for production environments. Packages that don't include the checkmark are often developed by other engineers or organizations and are not guaranteed to work in all situations. You can see an example of this in [Figure 3-10](#).

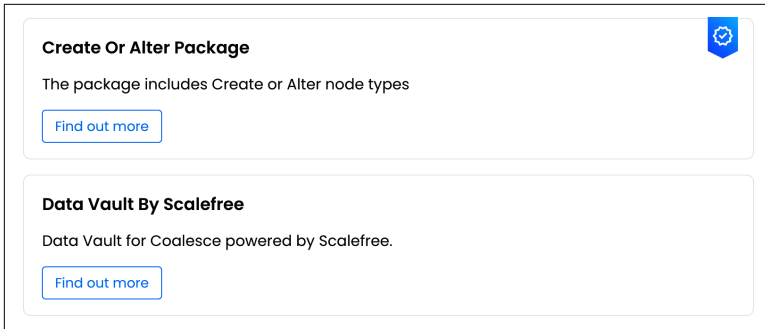


Figure 3-10. A package that has been certified by Coalesce (with the blue checkmark) and a package that has not been certified

You can easily install any of the node types from a package by clicking the “Find out more” button on any package. Within the package is a package ID, and this package ID is how the package is installed within Coalesce.

In the Build settings, you will see the Packages settings. Within the Packages settings, you will see the option to either browse or install packages, as seen in [Figure 3-11](#). The Browse button will open a new tab and take you to Coalesce Marketplace, where you can view all of the packages available to install. The Install button allows you to install a package by providing the package ID of the package from Coalesce Marketplace.

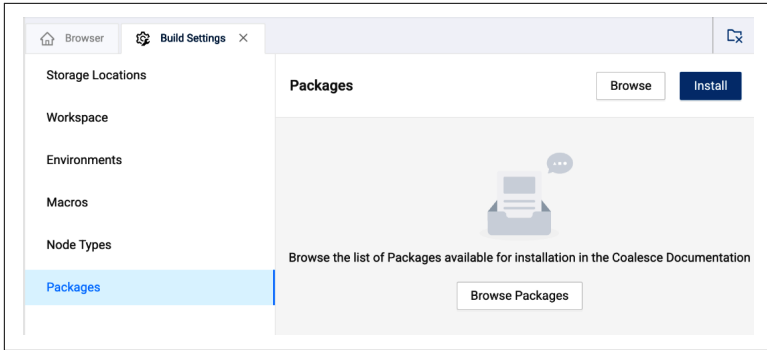


Figure 3-11. The Browse and Install buttons in the upper right corner of the Packages settings

When installing a package, you will provide a package ID, the version of the package, and an alias. By default, the package will install as the latest version, but you can manually change this to any previously supported version. The alias of the package is the name or alias under which the node types will be displayed in the Browser, as seen in Figures 3-12 and 3-13. Each node type that is installed from the package will be displayed in the Node Type settings within the Build settings of your Workspace. You can then add these node types to your pipeline just like any other node type.

Install Package

×

Packages are installed using a Package ID. For more details please view the Packages page in the [Coalesce Documentation](#).

Package ID

@coalesce/incremental-loading

Version

1.1.1 (latest) ▾

Package Alias

Incremental

Give your Package a unique and descriptive name. Learn more about [package aliases](#).

Install

Figure 3-12. Providing an alias when installing a package

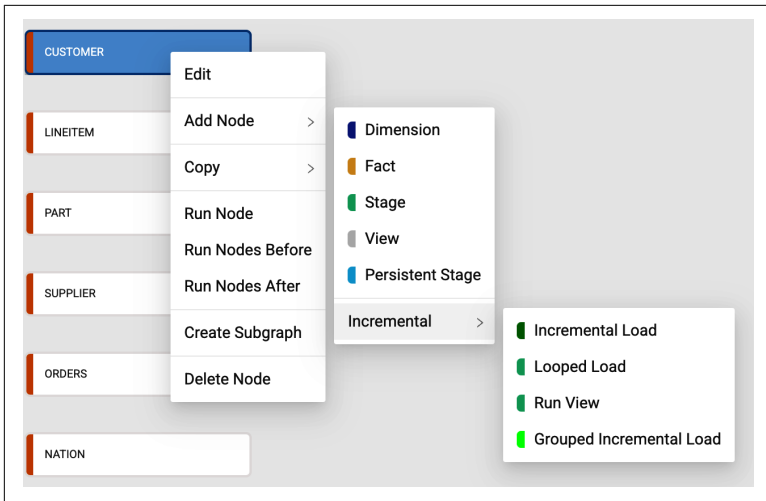


Figure 3-13. How the alias shows up when you add nodes in the Browser

Whether they are out of the box, from Coalesce Marketplace, or custom-built, you'll use the same method to add all node types to your data pipeline. Let's learn how to begin building on our data sources using different node types.

Putting Node Types to Work

As we discussed earlier in this chapter, a common pattern in data engineering is to create a staging layer to prepare your data for the needs of your business users. To do this in Coalesce, we can use the Stage node type. You can select any node, including a source node, in Coalesce and right-click on it and then hover over Add Node to view all of the node types available to add to your data pipeline. In this case, select Stage.

Coalesce will add a green Stage node type to the Browser and will automatically add a prefix to the node, `STG_`, as seen in [Figure 3-14](#). This naming convention provides an additional way to easily see the Stage nodes in your Browser.

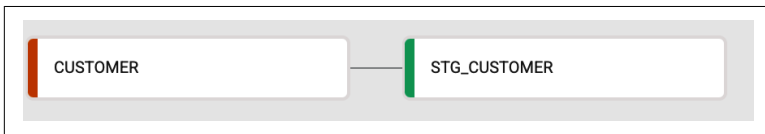


Figure 3-14. The Stage node type with the `STG_` prefix applied to the node

As you learned in [Chapter 2](#), Coalesce uses data patterns to provide the templates making up each node type. Because all node types are standardized objects, you can add them to your pipeline the same way. This also means that you can add them in bulk to multiple objects, since the standard never changes.

In the Browser, you can see multiple nodes selected at once. By right-clicking on any of the nodes selected, you can hover over Add Node and, to complete our staging layer, select Stage. Coalesce will automatically add a Stage node type to each data source, as seen in [Figure 3-15](#), allowing your team to immediately begin transforming your data without having to configure the code of each object individually.

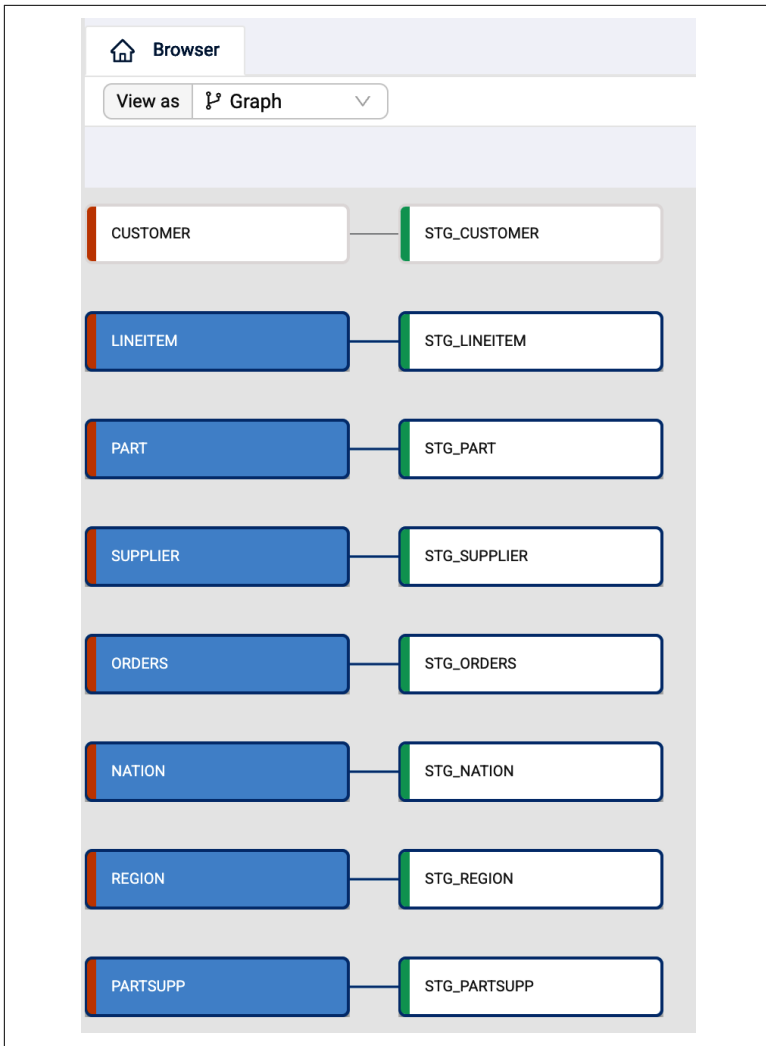


Figure 3-15. Bulk-adding Stage nodes to the data pipeline

You can add any other node type available in your Workspace the same way that you added the Stage node types. For instance, you can add a Dimension node by right-clicking on the `STG_CUSTOMER` node and selecting Dimension, as shown in [Figure 3-16](#).

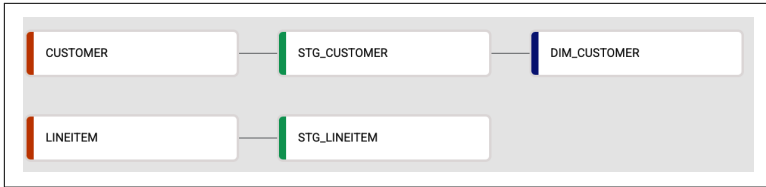


Figure 3-16. The `STG_CUSTOMER` node with a Dimension node as a dependency

You can continue this process for every object needed to develop your data in your pipeline. As you add nodes to your pipeline, you will need to configure them to help provide solutions to the problems you are solving. To do this, you will need to understand the anatomy of a node.

The Anatomy of a Node

In [Chapter 2](#), we discussed how a node type is created (Node Definition ► Create Template ► Run Template). In this chapter, you will learn about the basic anatomy of a node so you can effectively work with any node type. Whenever you double-click on a node in the Browser, it will open the Node Editor. While the configuration options (which you'll learn about shortly) may be different, the basic anatomy of each node is the same. Let's dive into this a bit deeper.

The mapping grid

When opening a node for the first time, you'll immediately see the mapping grid, as shown in [Figure 3-17](#). The mapping grid is the display of all of the columns, including the name, data type, transformations, and even comments that are inherited from the parent node(s) in the DAG. You can easily add new columns, apply transformations, change data types, and apply a multitude of other operations in the mapping grid.

Column Name	Transform	Data Type	Source	Nullable	Description
C_CUSTKEY		NUMBER(38,0)	"CUSTOMER"."C_CUSTKEY"	false	
C_NAME		VARCHAR(25)	"CUSTOMER"."C_NAME"	false	
C_ADDRESS		VARCHAR(40)	"CUSTOMER"."C_ADDRESS"	false	
C_NATIONKEY		NUMBER(38,0)	"CUSTOMER"."C_NATIONKEY"	false	
C_PHONE		VARCHAR(15)	"CUSTOMER"."C_PHONE"	false	
C_ACCTBAL		NUMBER(12,2)	"CUSTOMER"."C_ACCTBAL"	false	
C_MKTSEGMENT		VARCHAR(10)	"CUSTOMER"."C_MKTSEGMENT"	true	
C_COMMENT		VARCHAR(117)	"CUSTOMER"."C_COMMENT"	true	

Column Name Transform Data Type Source Nullable Description

Figure 3-17. The mapping grid of the `STG_CUSTOMER` node

Every node in your DAG will contain a mapping grid, even if it includes only a single column or line item, and it will tell you which columns you are working with.

The configuration options

Each node type contains configuration options that are unique to that node type. For example, the Dimension node type contains configuration options such as a business key and change tracking columns. These are different from the options for the Stage node type, for example. The configuration options of each node type allow users to accelerate their data development by reusing what has already been created as a standard. This is why configuration of a type 2 SCD saves hours for data developers in Coalesce, because you can configure with just a few clicks while it automatically generates hundreds of lines of SQL.

You can view the configuration options of any node in the Config tab in the upper right corner of the Node Editor, as seen in [Figure 3-18](#).

The screenshot shows the 'Config' tab of a node configuration window. At the top, there are three tabs: 'Config' (selected), 'Column Editor', and 'Testing'. Below the tabs is a sidebar with two items: 'Node Properties' (expanded) and 'Options' (collapsed). The main area contains the following settings:

- Create As ***: A dropdown menu with 'table' selected.
- Multi Source**: A toggle switch that is currently turned off.
- Truncate Before**: A toggle switch that is currently turned on.
- Enable Tests**: A toggle switch that is currently turned on.
- Pre-SQL**: A section with an 'Expand' button and a text area containing the number '1'.
- Post-SQL**: A section with an 'Expand' button and a text area containing the number '1'.

Figure 3-18. The configuration options, providing the ability to configure a node without having to write code

Create and run

A critical part of developing data pipelines in Coalesce is the ability to create and run nodes. But what does that mean? When you add a node to your pipeline in Coalesce, the node is not immediately created in your data platform, e.g., Snowflake. You as the data developer need to create the object in your data platform. The Create

and Run buttons allow you to create an object and then execute any action on the object. Let's break this down in an example.

When working with a Stage node, you have the ability to materialize the object as either a table or a view. Let's assume you choose to build your object as a table. In order for the object to be created in your data platform, you need to select the Create button. When you do so, Coalesce will automatically use all of the metadata about the node to generate the DDL for the object and create a blank object in your data platform. Coalesce will use the name of the node as the name of the object in your data platform, as shown in Figures 3-19 and 3-20.

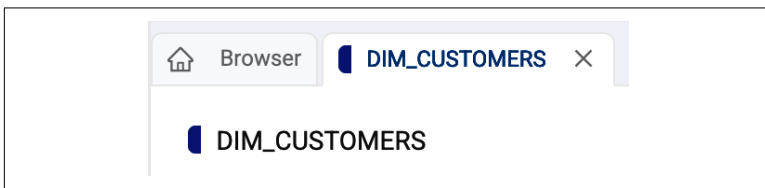


Figure 3-19. Name of the object in Coalesce

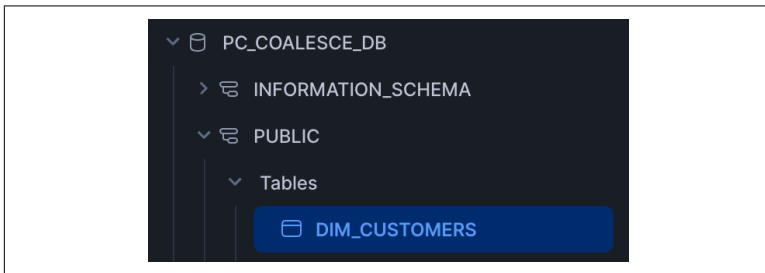


Figure 3-20. The same object created with the same name in your data platform

Once the Stage is created as a table in your data platform, it is still empty. This is where the Run button comes in. By selecting Run within the node, you will again tell Coalesce to use all of the metadata provided from the node (column names, transformations, configurations, etc.) to automatically generate the DML for the object, inserting data into the table. You'll be able to see the output of the operation in the Data Preview pane, which will display any data available after a successful run.

This is how Coalesce acts as the interface between your data platform and your raw data. Now you may be thinking that manually

creating and running each node is not a particularly sustainable approach to pipeline development, and you would be right! In [Chapter 4](#), you will learn about deployments and how you can automatically schedule your data pipelines to refresh, which effectively runs any of the Create and Run operations as needed.

The Join tab

Each node type contains a Join tab. This is located next to the mapping grid, as shown in [Figure 3-21](#). Within the Join tab, you will notice that there is a line of SQL already in the editor. The SQL shows a FROM statement with a REF function nestled within some curly braces. This line of SQL is creating the dependency to the upstream or parent node, i.e., the REFerence to the parent.

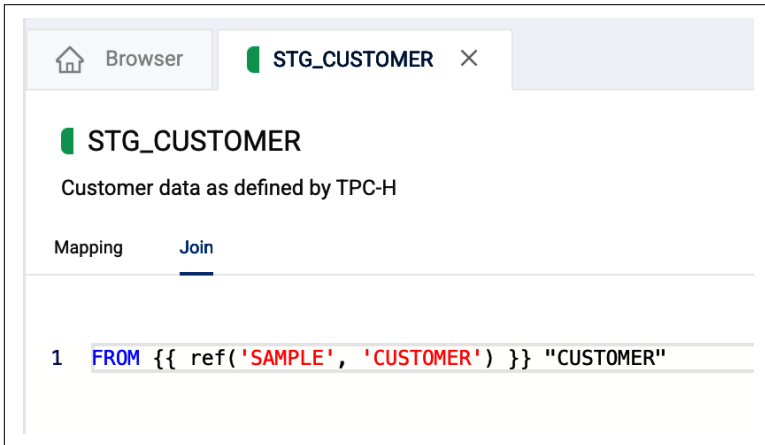


Figure 3-21. The Join tab next to the mapping grid

Within the Join tab, you can perform multiple operations such as joins, filters, and even window functions. The Join tab is often used for applying transformations to the object as a whole, rather than to a singular column. We'll dig more into the Join tab shortly, but for now, you only need to know where it is and what it can do.

You may notice a few other options within each node, and we'll get to those later in the chapter, but now it's time to get to the core of what nodes are doing: transforming data!

Data Transformations in Coalesce

In this section, I will provide a general overview of core data transformation functionality in Coalesce, but know that there are virtually an unlimited number of ways you can transform your data in Coalesce using various node types and settings. To cover the basics, I'll focus on column-level and node-level transformations.

Column-Level Transformations

To kick things off, let's start with *column-level transformations*. These are applied to individual columns within the mapping grid and are written using any valid SQL expression supported by your cloud platform. Let's look at an example.

In the node in [Figure 3-22](#), we are applying an UPPER() function to a variable character column in order to apply consistent text casing across the field.

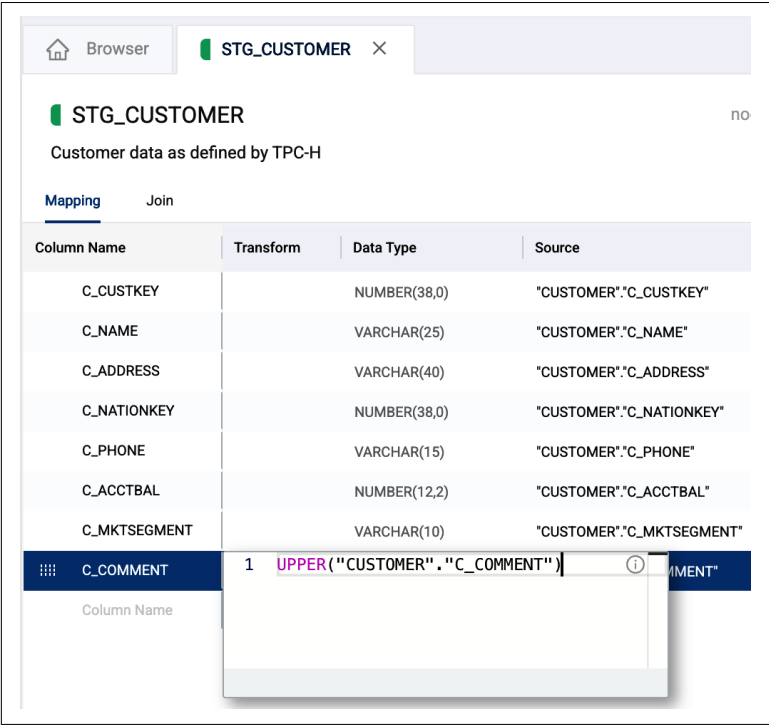
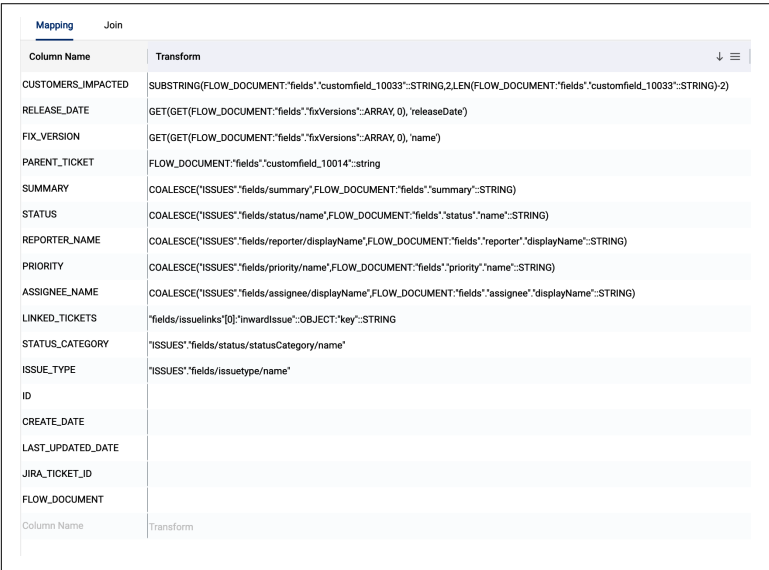


Figure 3-22. Applying a column-level transformation to a column in the mapping grid

While this is a relatively simple example of a column-level transformation, you can probably begin to see the ease of managing transformations in this way. For another example, take a look at the column-level transformations shown in [Figure 3-23](#).



The screenshot shows the 'Mapping' tab in the Coalesce interface. It displays a table with two columns: 'Column Name' and 'Transform'. The table lists various columns and their corresponding transformation expressions. The 'Transform' column contains complex expressions using Coalesce's syntax, such as `SUBSTRING(FLOW_DOCUMENT::fields::customfield_10033::STRING,2,LEN(FLOW_DOCUMENT::fields::customfield_10033::STRING)-2)` for `CUSTOMERS_IMPACTED` and `COALESCE('ISSUES::fields/summary',FLOW_DOCUMENT::fields::summary::STRING)` for `SUMMARY`. The table is scrollable, and a summary row at the bottom shows the column headers.

Column Name	Transform
CUSTOMERS_IMPACTED	<code>SUBSTRING(FLOW_DOCUMENT::fields::customfield_10033::STRING,2,LEN(FLOW_DOCUMENT::fields::customfield_10033::STRING)-2)</code>
RELEASE_DATE	<code>GET(GET(FLOW_DOCUMENT::fields::fixVersions::ARRAY,0),releaseDate)</code>
FIX_VERSION	<code>GET(GET(FLOW_DOCUMENT::fields::fixVersions::ARRAY,0),name)</code>
PARENT_TICKET	<code>FLOW_DOCUMENT::fields::customfield_10014::string</code>
SUMMARY	<code>COALESCE('ISSUES::fields/summary',FLOW_DOCUMENT::fields::summary::STRING)</code>
STATUS	<code>COALESCE('ISSUES::fields/status/name',FLOW_DOCUMENT::fields::status::name::STRING)</code>
REPORTER_NAME	<code>COALESCE('ISSUES::fields/reporter/displayName',FLOW_DOCUMENT::fields::reporter::displayName::STRING)</code>
PRIORITY	<code>COALESCE('ISSUES::fields/priority/name',FLOW_DOCUMENT::fields::priority::name::STRING)</code>
ASSIGNEE_NAME	<code>COALESCE('ISSUES::fields/assignee/displayName',FLOW_DOCUMENT::fields::assignee::displayName::STRING)</code>
LINKED_TICKETS	<code>'fields/issuelinks[0]:inwardIssue':OBJECT::key::STRING</code>
STATUS_CATEGORY	<code>'ISSUES::fields/status/statusCategory/name'</code>
ISSUE_TYPE	<code>'ISSUES::fields/issuetype/name'</code>
ID	
CREATE_DATE	
LAST_UPDATED_DATE	
JIRA_TICKET_ID	
FLOW_DOCUMENT	
Column Name	Transform

Figure 3-23. Multiple column transformations easily managed in the mapping grid

Column-level transformations can be quite complex, such as when they use nested CASE WHEN statements that contain aggregate functions. Effectively, if you can write the column-level transformation in your data platform, you'll be able to write it within Coalesce.

Each column-level transformation will contain the name of the column being transformed, as well as the upstream dependency of that column. In this way, you are listing the full table and column reference to provide the metadata that will generate the transformation code automatically.

It's important to note that column-level transformations only transform singular columns. You will not be able to filter an entire table using a column-level transformation. That's where you would use the Join tab, which we discuss next.

Node-Level Transformations

Using the Join tab, you can apply table- or node-level transformations to your data. Let's assume you've transformed all of your columns and are ready to apply any SQL necessary to support the logic of the table. For example, maybe you used an aggregate function and now need to supply a `GROUP BY`; you can supply this in the Join tab. Or maybe you're working with ecommerce sales data, as shown in [Figure 3-24](#), and only want to view sales where the order was delivered; you could apply the filter in the Join tab. Effectively, if you need to apply any SQL at the table level, you should apply it in the Join tab.

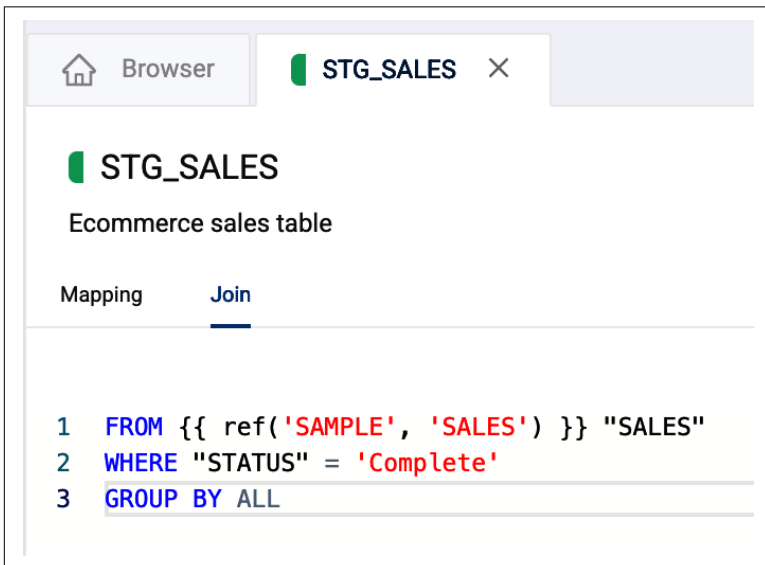


Figure 3-24. Writing SQL to configure the node at the node level

Whether you are using column-level transformations or transforming data at the node level, you can use any SQL supported by your data platform to do so. There is no proprietary syntax limiting your ability to write SQL or use a different language, so you can just plug and play. Now that you're familiar with the Join tab, let's explore its namesake—the join itself!

Joins

Up until this point, we have been working with and transforming data in single nodes. But what happens when you want to join multiple nodes together? Coalesce makes this just as easy as adding a node to your data pipeline.

In the Browser, you can select two or more nodes you want to join together. Once they're selected, you can right-click on either node and hover over Join Nodes, and you will see all of the options available for adding a node but now available for a joined node. In [Figure 3-25](#), the STG_ORDERS and STG_LINEITEM nodes can be joined to make a new Stage node type. [Figure 3-26](#) shows the result of a join in the Browser.

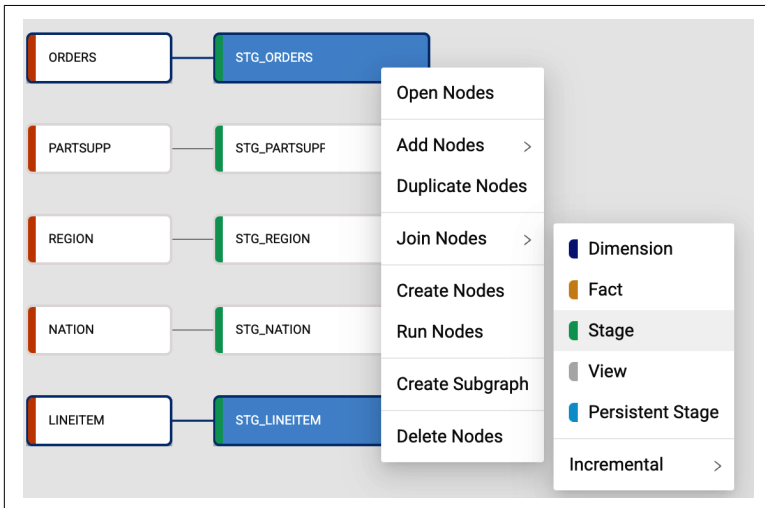


Figure 3-25. Applying a join using the Join Nodes option in the Browser

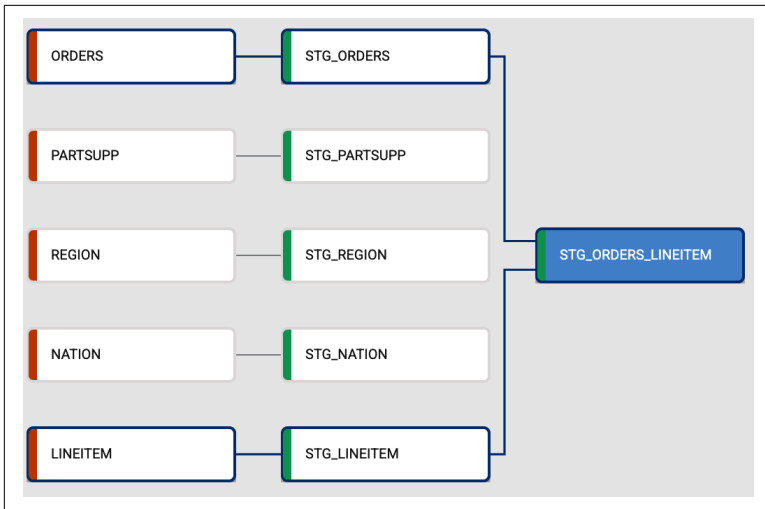


Figure 3-26. The result of using the Join Nodes option in the Browser

Once the joined Stage node is selected, Coalesce will automatically drop you into the Node Editor. To configure the join, navigate to the Join tab. You will see a join statement that Coalesce has automatically generated for you. In most cases, all you need to do is provide the join condition (i.e., the column names used to join the tables together). This can be done by removing the `/*COLUMN*/` placeholders (Figure 3-27) and replacing them with the column names needed to configure the Join (Figure 3-28).

Mapping	Join
<pre> 1 FROM {{ ref('WORK', 'STG_ORDERS') }} "STG_ORDERS" 2 INNER JOIN {{ ref('WORK', 'STG_LINEITEM') }} "STG_LINEITEM" 3 ON "STG_ORDERS"./*COLUMN*/ = "STG_LINEITEM"./*COLUMN*/ </pre>	

Figure 3-27. Column placeholders in the Join tab for join conditions

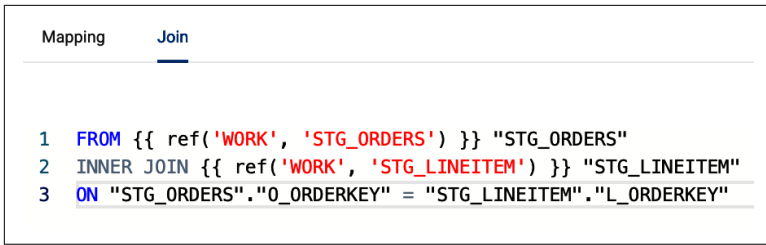


Figure 3-28. Configuration of the join in the table with the proper column condition

As you learned earlier, when it comes to table-level transformations, we can always add another join condition if necessary, as shown in Figure 3-29.

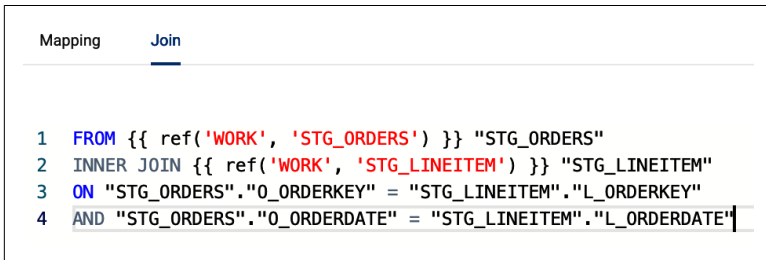


Figure 3-29. Adding another join condition, showing the flexibility of the SQL editor

Coalesce automatically infers when a join is occurring. If you happen to delete or break your code, you can always have Coalesce regenerate the join. In the upper right corner of the Join tab SQL editor, you'll see the Generate Join dropdown (Figure 3-30), which allows you to regenerate the code in the SQL editor or copy it directly to your clipboard.

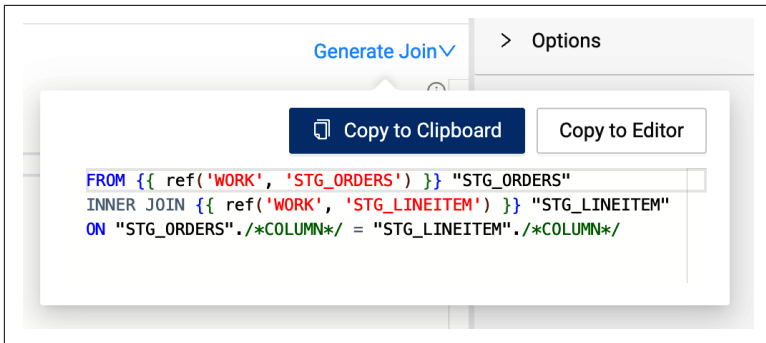


Figure 3-30. The automatic join generation that Coalesce provides in each node

This same process can be used to join any number of nodes together, making it management of even the most complex joins simple and easy.

Once tables have been joined, there are often many columns that need additional work—whether it is renaming, changing data types, or applying transformations. In Coalesce, you can handle these tasks efficiently using the bulk editor, which you will learn about next.

Bulk Editing

Joins can mean you have to deal with multiple columns at once. In traditional data development scenarios, you may need to manually delete columns one by one or apply transformations to one column at a time. But because Coalesce uses data patterns and metadata to accelerate your data development, you can bulk-edit columns straight from the mapping grid.

Let's assume you've joined the two nodes together from our previous example in [Figure 3-28](#). You now have multiple columns containing variable character datatypes. In the case where you want to apply consistent text casing, as in [Figure 3-22](#) earlier, you can bulk-add this transformation to each column you select and then write the transformation only once.

If you were to select all of the VARCHAR columns in the node and right-click on any of the columns, you could select Bulk Edit, as shown in [Figure 3-31](#). This would open the column editor next to the configuration options that you learned about at the beginning of this chapter.

Column Name	Transform	Data Type	Source	Nullable	Description
L_QUANTITY		NUMBER(12,2)	"STG_LINEITEM"."L_QUANTITY"	false	
L_EXTENDEDPRICE		NUMBER(12,2)	"STG_LINEITEM"."L_EXTENDEDPRICE"	false	
L_DISCOUNT		NUMBER(12,2)	"STG_LINEITEM"."L_DISCOUNT"	false	
L_TAX		NUMBER(12,2)	"STG_LINEITEM"."L_TAX"	false	
O_ORDERKEY		NUMBER(38,0)	"STG_ORDERS"."O_ORDERKEY"	false	
O_CUSTKEY		NUMBER(38,0)	"STG_ORDERS"."O_CUSTKEY"	false	
O_SHIPPRIORITY		NUMBER(38,0)	"STG_ORDERS"."O_SHIPPRIORITY"	false	
L_ORDERKEY		NUMBER(38,0)	"STG_LINEITEM"."L_ORDERKEY"	false	
L_PARTKEY		NUMBER(38,0)	"STG_LINEITEM"."L_PARTKEY"	false	
L_SUPPKEY		NUMBER(38,0)	"STG_LINEITEM"."L_SUPPKEY"	false	
L_LINENUMBER		NUMBER(38,0)	"STG_LINEITEM"."L_LINENUMBER"	false	
O_ORDERSTATUS		VARCHAR(1)	"STG_ORDERS"."O_ORDERSTATUS"	false	
L_RETURNFLAG		VARCHAR(1)	"STG_LINEITEM"."L_RETURNFLAG"	false	
L_LINESTATUS		VARCHAR(1)	"STG_LINEITEM"."L_LINESTATUS"	false	
L_SHIPMODE		VARCHAR(10)	"STG_LINEITEM"."L_SHIPMODE"	false	
O_ORDERPRIORITY		VARCHAR(15)	"STG_ORDERS"."O_ORDERPRIORITY"	false	
O_CLERK		VARCHAR(15)	"STG_ORDERS"."O_CLERK"	false	
L_SHIPINSTRUCT		VARCHAR(25)	"STG_LINEITEM"."L_SHIPINSTRUCT"	false	
L_COMMENT		VARCHAR(44)	"STG_LINEITEM"."L_COMMENT"	false	
O_COMMENT		VARCHAR(79)	"STG_ORDERS"."O_COMMENT"	false	

Bulk Edit
 Duplicate Columns
 Add Node >
 View Column Lineage
 Generate Hash Column
 Delete Columns

Column Name Data Type Source Nullable Description

Selected: 9 Rows: 26

Figure 3-31. Selecting columns in bulk in order to bulk-edit them

You could then select the attribute you wanted to transform, such as the column names, data types, or transformations. For our example, we'll write a simple `UPPER({{SRC}})` function that can be applied to each column, as shown in [Figure 3-32](#). In this case, the `{{SRC}}` automatically resolves the column in any transformation!

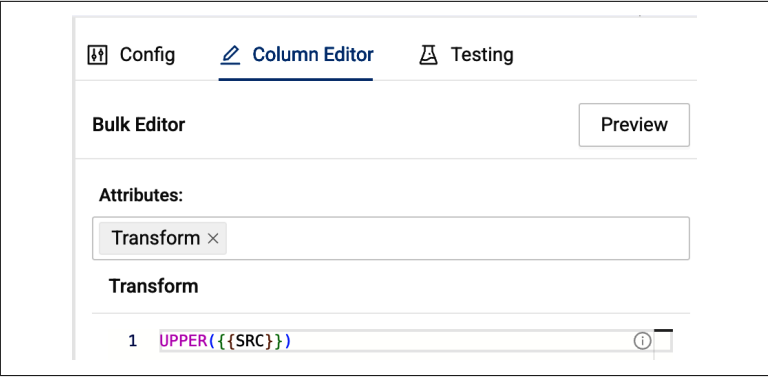


Figure 3-32. The bulk column editor, applying the *UPPER* function to multiple columns as a transformation

An attribute of a column can be bulk-transformed in this way. Additionally, if we wanted to remove multiple columns, we could easily select all of the columns we wanted to delete, right-click on any of the selected columns, and select **Delete Columns**, as shown in [Figure 3-33](#).

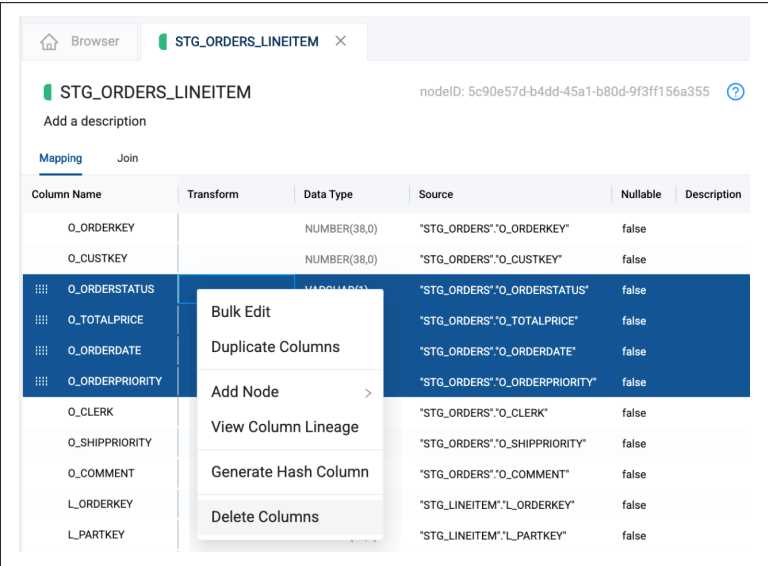


Figure 3-33. Bulk-deleting columns from the mapping grid

By allowing users to work with columns in bulk, Coalesce makes it easy to update multiple fields at once—saving time and letting developers focus on higher-value work.

Data Transformation in Process

There was quite a lot of information to digest in this chapter. You learned about the build interface and all of the components it contains. You learned all about adding data sources and nodes to your pipeline, including the many kinds of node types available to you. You also learned about the anatomy of a node and how to use SQL to transform data inside of any node, while also joining nodes together and applying bulk updates. Whew!

Coalesce is no less powerful because it leads with a user interface. In fact, it enhances efficiency by standardizing the way work is done, ensuring that everyone builds from the same foundation. And when you need to write SQL, it is always available. A visual interface does not limit capability—it amplifies it.

In the next chapter, you'll sharpen your new data transformation skills by learning how to manage the data pipelines you build in Coalesce. You're quickly becoming a data development master.

Managing Data Pipelines in Coalesce

In the previous chapter, you used Coalesce to build data transformations and pipelines. But managing existing pipelines is just as crucial as creating new ones. This chapter covers both: how to build any data product in Coalesce and how to manage pipelines of any size effectively.

You'll learn how to inspect objects to understand their contents and where they are built. You'll also use column-level lineage to perform impact analysis and leverage the Problem Scanner to identify and resolve issues efficiently. This chapter also explores key management features, including version control, macros and parameters, testing, subgraphs, jobs, and environments.

First, you'll navigate the nodes in your pipeline to establish a solid foundation for managing your data workflows.

Managing Nodes Using Views

As you build nodes in Coalesce, obtaining full context about them is often essential for development decisions. Fortunately, Coalesce makes this process straightforward. By default, it displays pipeline development as a graph, as shown in [Figure 4-1](#).

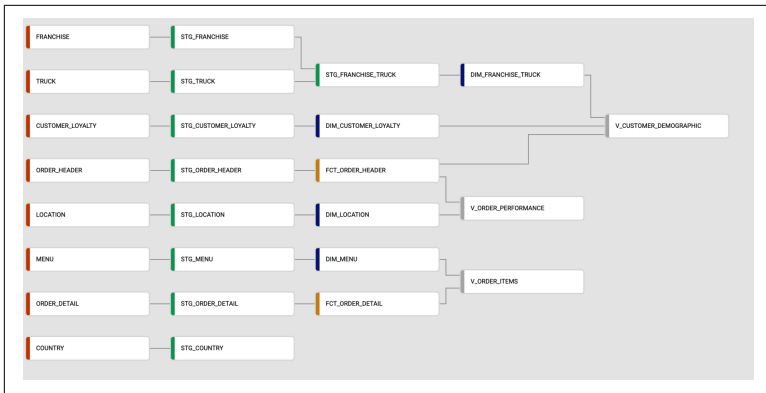


Figure 4-1. Multiple node types displayed as a graph in Coalesce

Coalesce provides two more views as well: the node grid and the column grid. These views offer deeper insights into your pipeline. Let's explore all three views.

The Graph

The *graph view* provides the best way to see your entire pipeline and its dependencies. It constructs a DAG that displays the connections and dependencies of each node. This view gives a comprehensive understanding of how data moves through the pipeline. When you need to troubleshoot, the graph view helps you quickly identify dependencies and determine how far upstream or downstream a specific object is. It also makes it easy to distinguish nodes by their unique colors, as shown in [Figure 4-2](#).

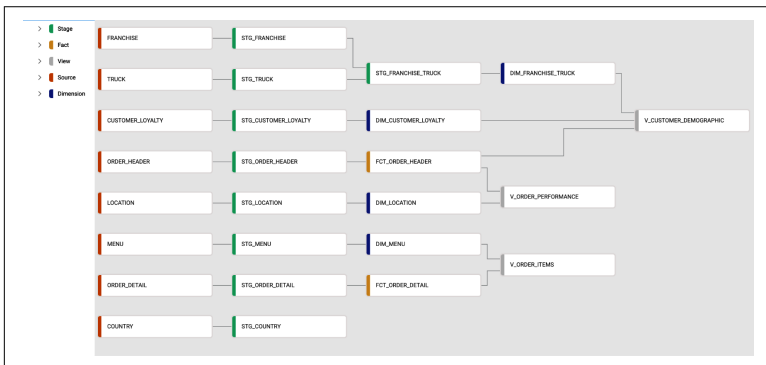


Figure 4-2. The node menu rolling up all of the nodes to the respective color of the node type

The Node Grid

When developing data products in Coalesce, finding a specific node or object in the graph can be challenging, especially in projects with hundreds or thousands of nodes. For large projects, you can use subgraphs to help with this, and you’ll learn more about that later in “Subgraphs” on page 70, but right now let’s learn about the node grid.

The *node grid* provides several ways to filter and locate nodes efficiently. It allows filtering by individual columns, so you can display only nodes of a specific type, such as Fact nodes. Additionally, the filter bar supports selector queries, which you’ll learn about in “Selector Queries” on page 72, enabling you to include or exclude values from your search to refine the displayed nodes. For example, if you are processing orders data with some Stage nodes, you can use the filter bar and selector syntax to search for nodes that follow a STG_ORDER naming convention, as shown in Figure 4-3.

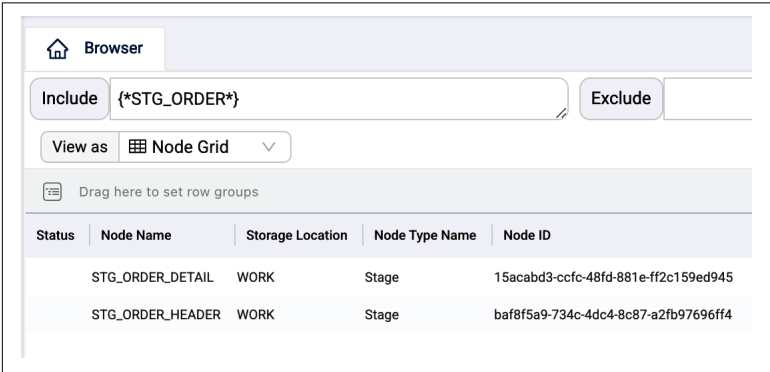


Figure 4-3. Using an Include filter to filter the node grid for only tables with STG_ORDER in the name

You can also group nodes in the mapping grid by their associated storage location. To do this, drag and drop the Storage Location column into the row groups area. This immediately organizes nodes into groups based on their storage location. You can further refine these groups using column filters and selector queries.

The Column Grid

Column-level architecture, which you have already explored in previous chapters, makes it easy to manage columns in Coalesce. A primary way to do this is through the *column grid*, which provides a detailed view of every column in your workspace in one place. This view offers several advantages.

First, the column grid displays all columns in a single location. Therefore, this view provides complete transparency. It allows data practitioners to see all the columns in a workspace without searching through individual objects or fragmented documentation.

Second, the column grid supports the same filtering and selector queries as the node grid. This makes it easy to refine searches and locate specific columns. For example, you can apply a filter to find any column containing a specific name, as shown in [Figure 4-4](#).

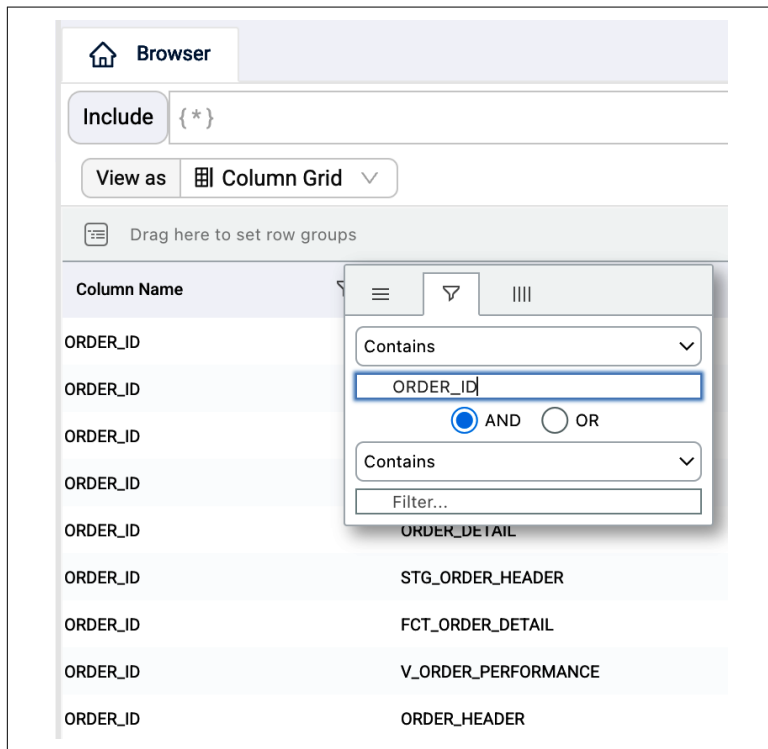


Figure 4-4. Using a column header filter in the column grid to filter for `ORDER_ID`

Like the node grid, the column grid allows you to create groupings using column headers. For example, you can group by “Storage Location” and “Node Name” to refine searches for columns within a specific location or object in your data platform.

You can also use selector queries or column filtering to view and manage transformations applied to any column. This provides a powerful interface for handling column-level transformations. For example, you can search for a specific column name and review all of that column’s associated transformations. This is especially useful when updating a transformation across multiple objects.

For instance, let’s assume you use a transformation to extract the area code from phone number columns in various datasets. This transformation relies on a LEFT function to capture the first three digits of the phone number. However, after an application update, phone numbers now include a country code (such as +1 for the United States). As a result, the existing transformation no longer works correctly because it captures the country code along with the first one or two digits of the area code.

To update this transformation, you can filter the mapping grid to locate the specific transformation. Then, using the bulk editor, you can update every affected column at once by modifying the function and applying the changes with a single click, as shown in [Figure 4-5](#).

Column Name	Node Name	Storage Location	Transform	Source
workspace_id	STG_ORGS_WORKSPACES_STEPS	STAGING	SPLIT_PART(meta_path/, 9)	
org_id	STG_ORGS_WORKSPACES_STEPS	STAGING	SPLIT_PART(meta_path/, 7)	
node_id	STG_ORGS_WORKSPACES_STEPS	STAGING	SPLIT_PART(meta_path/, 11)	
refresh_success_rate	FCT_MONTHLY_REFRESH_SUCCESS	TARGET	NVL(STG_COMPLETED_REFRESH_RU	
deploy_success_rate	FCT_MONTHLY_DEPLOY_SUCCESS	TARGET	NVL(STG_COMPLETED_DEPLOY_RUN	
csat_score	STG_MONTHLY_CSAT	STAGING	NVL(SUM(CASE WHEN 'STG_ACCOUNT	
first_month	STG_FIRST_MONTH	STAGING	MIN(DATE_TRUNC('month', STG_ARRL	'STG_ARRL_OPPS', 'date_closed'

Figure 4-5. Bulk-editing columns after identifying the transformations in the column grid

Each view in Coalesce offers distinct ways to manage objects and columns in your data pipeline. These views equip you with the tools needed to efficiently navigate and control your pipelines. However, Coalesce provides additional functionality that extends beyond node visualization, allowing for even greater project management control.

Subgraphs

So far, you have explored how to use Coalesce pipeline views to manage objects in your pipeline. As your workspace grows to hundreds or thousands of nodes, viewing everything in the graph can become challenging. *Subgraphs* help manage this complexity by allowing you to isolate parts of the node graph.

Subgraphs break down defined sections of your pipeline into logically separate sections. They can be defined manually or by using a selector query. For example, suppose you are working on an enterprise data warehouse in Coalesce with dozens of data sources. Your graph has grown to over five hundred nodes, making it difficult to locate nodes associated with a specific data source, such as customer relationship management (CRM) data. To address this, you can organize each data source and its processing nodes into separate subgraphs.

To create a subgraph, select all nodes related to your CRM data that you want to associate with the subgraph. Then, in the Subgraphs tab, click the plus (+) button and select Create Subgraph. Coalesce will generate a subgraph that displays only the nodes associated with your CRM data, as shown in [Figure 4-6](#).

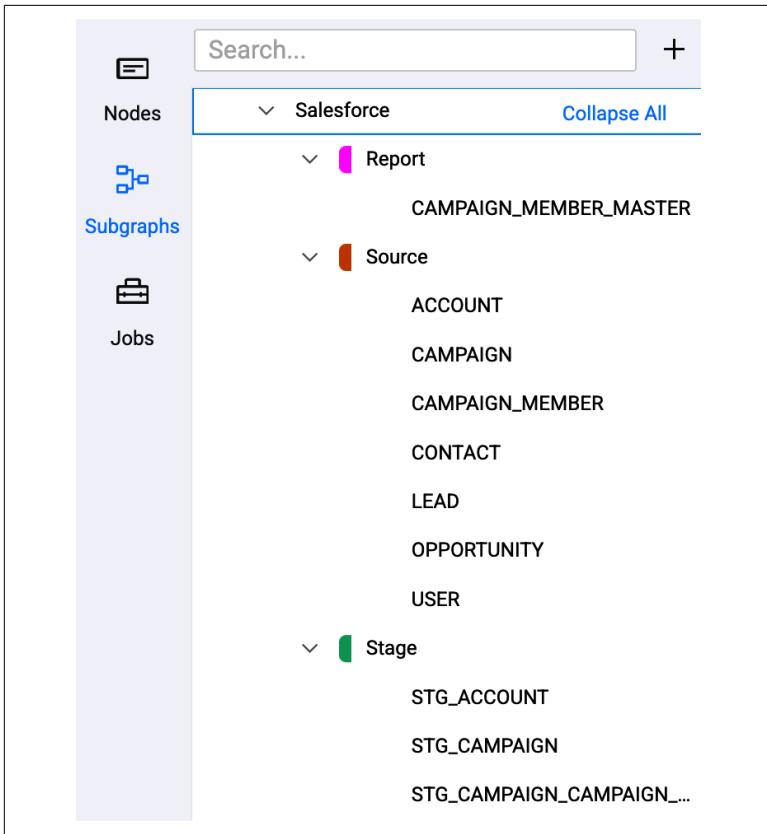


Figure 4-6. Subgraph created to contain only data from a CRM (Salesforce)

Additionally, just as you can use selector syntax for viewing nodes, you can use it to select nodes to add to a subgraph, as shown in [Figure 4-7](#).

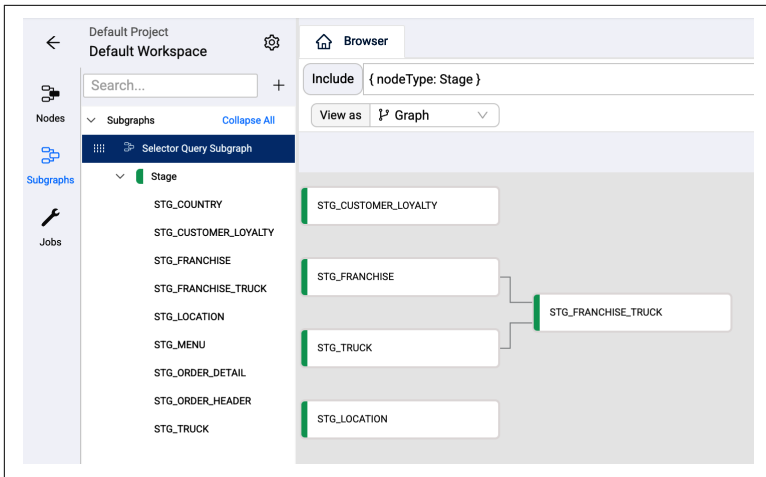


Figure 4-7. Using selector syntax to select the nodes to add to a subgraph

Up to this point, you have seen multiple references to selector syntax, but you have not yet explored what it is. Now that you know where you can use selector syntax, let's take a closer look at what it is and how to apply it.

Selector Queries

Selector queries refine and enhance search criteria throughout various functions in Coalesce. Each selector query follows a structured search syntax with attributes that allow you to retrieve specific information.

For example, to filter your pipeline for only Stage nodes, you can use the `nodeType` attribute in the following query:

```
{ nodeType: Stage }
```

This query filters the workspace to display only Stage nodes. You can apply the same approach to any other node type.

Coalesce supports multiple attributes, which are documented in the [Coalesce documentation](#), to help refine and enhance queries. You can combine multiple attributes within a single query—for example:

```
{name: SUPPLIER* location: BI nodeType: Dimension }
```

This query searches for nodes that meet all three of the following criteria:

- The node name contains SUPPLIER (the asterisk [*] is used as a wildcard).
- The node is built within the BI storage location.
- The node is a Dimension node.

Operators further enhance selector queries by providing more control over searches. For example, the plus (+) operator selects predecessors or successors based on placement:

```
+{name:STG_SUPPLIER}  
{name:STG_SUPPLIER}+
```

You can also use the AND and OR operators to add conditional logic:

```
**{name:STG_} OR {name:REGION}
```

Selector queries allow precise filtering of pipelines, giving you more control over how your data is managed. You will revisit selector queries in “Jobs” on page 86. But for now, you should feel comfortable using them to filter views and create subgraphs. Selector queries aren’t the only way to refine your searches within your data pipeline. Coalesce also provides column-level lineage to allow you to understand how each column flows through your pipeline, which you’ll learn about next.

Column-Level Lineage

In [Chapter 2](#), you explored column-level architecture in Coalesce. Since Coalesce is built with column-level architecture from the ground up, you can take advantage of powerful features that streamline data pipeline management and accelerate development. One of the most powerful components that results from this is column-level lineage.

Column-level lineage in Coalesce allows you to track each column as it moves through the pipeline. This includes transformed columns, showing how multiple columns combine into one or how a single column splits into multiple columns. Leveraging its column-level architecture, Coalesce also enables you to propagate column additions and deletions to downstream nodes with a single click.

To view column-level lineage, double-click into any node, select one or more columns, right-click the selection, and choose View Column Lineage. Coalesce will open the column lineage view, highlighting the selected column or columns and displaying their flow throughout the entire pipeline, as shown in **Figure 4-8**.



Figure 4-8. Column-level lineage view showing how columns flow through a pipeline and are transformed

This context enables you to perform quick and precise impact analysis, since it allows you to see exactly how changes will affect your pipeline. Additionally, any column with a transformation can display its transformation within column-level lineage, providing further insight into how modifications will impact your pipeline, as shown in **Figure 4-9**.

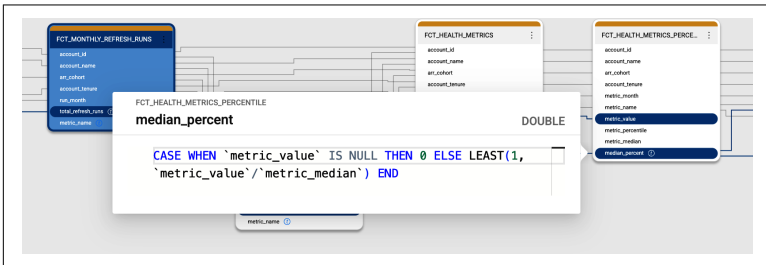


Figure 4-9. Column-level lineage displaying the transformation involved in the column

For example, suppose your extract, transform, load (ETL) solution adds a new column that needs to be exposed in a dashboard for stakeholders. In traditional pipeline development, this process can be tedious, requiring you to open every object the column passes through, run new DDL scripts, and potentially perform complex backfilling operations. In Coalesce, you can propagate the new column to all necessary nodes in just a few clicks.

In this scenario, the pipeline is undergoing a schema change. Let's assume the pipeline processes data through a staging layer before exposing it to downstream nodes. You can rename or transform the

It's really that simple. Coalesce handles all necessary DDL changes, updates each selected node with the new column, and allows you to process data as needed within your nodes.

Column-level lineage provides a clear view of how each column moves through your data pipelines. Additionally, it enables you to propagate changes with just a few clicks, reducing errors and accelerating development.

Version Control

Version control is a critical part of any data development project. Coalesce provides robust version control, enabling a continuous integration and continuous development (CI/CD) process through your preferred Git provider. In this guide, we assume you are using GitHub.

In [Chapter 1](#), you learned that Coalesce natively integrates with Git. Now it is time to understand how version control works within Coalesce. Once Git is integrated, you can create a project and a workspace. The project links to the Git repository of your choice, while the workspace is associated with a branch from that repository, as shown in [Figure 4-12](#).



Figure 4-12. Workspace with a Git branch tied to it

With your repository and branch configured, you can launch your workspace and begin building data products as you've learned to do in the past two chapters. As you build, you will notice the Git icon, which shows branching with nodes, in the left navigation bar toward the lower left corner, as shown in [Figure 4-13](#).

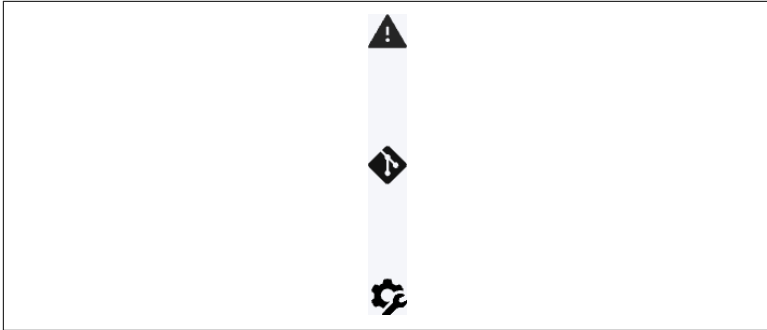


Figure 4-13. The Git icon in the middle of the navigation bar

When you click on the Git icon, the Coalesce Git modal will open, and you will be able to view all of the work you've performed in your workspace since your last commit. If you've never performed a commit, the modal will display all of your work, broken down into the categories the files are associated with, such as nodes, subgraphs, and packages, as shown in [Figure 4-14](#).

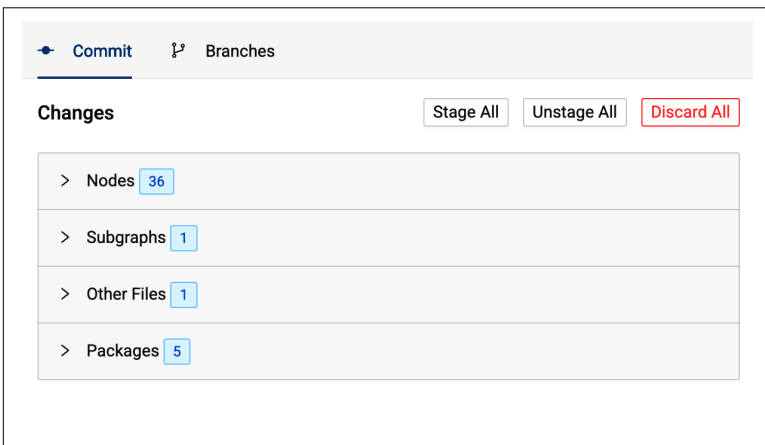


Figure 4-14. Categories of files, ready for version control

When you open any version-controlled file, you can view a *file differential (diff)*. For new files, the diff will be blank. For recently changed files, Coalesce highlights modifications using red and green lines, as shown in [Figure 4-15](#). In Coalesce, all work is saved as YAML files, so diffs are displayed in YAML format.



Figure 4-15. Coalesce highlighting the file diffs in the Git modal

Once you have reviewed your diffs and approved your work, you can choose to commit part or all of your work at once. Within the Git modal, by default, every file is selected to be included in a commit. You can deselect any of the files you wish to exclude by clicking on the blue checkbox next to the filename, as seen in [Figure 4-16](#).

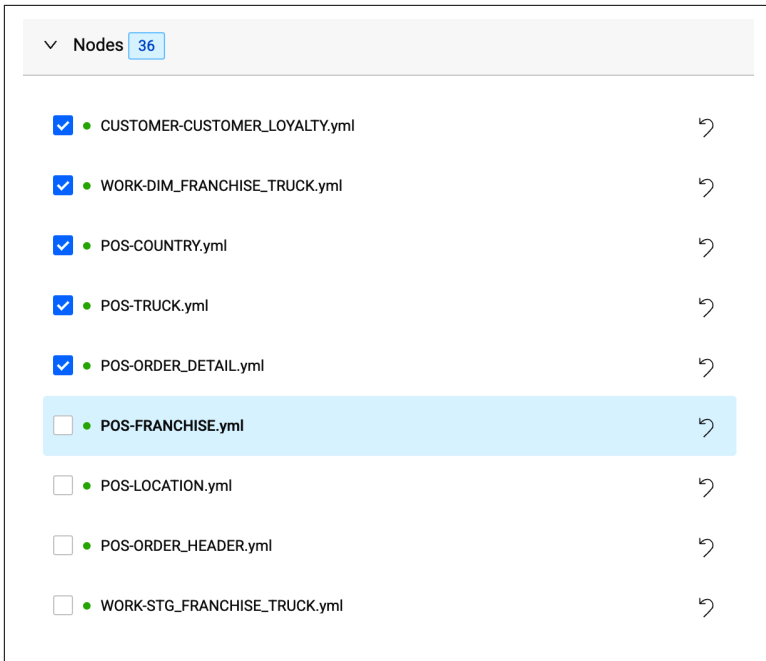


Figure 4-16. Selecting only certain files to be committed

After selecting the files to commit, you can write a commit message or use the Coalesce AI commit message generator to generate a message based on your changes. Once the commit message is in place and the files are selected, click Commit and Push to commit your changes and push them directly to your version control system from Coalesce.

So far, you have explored the Commit section within the Git modal. However, you can also check out, merge, or create a new branch from any commit within the modal. To do this, select the Branches tab next to the Commit tab, as shown in [Figure 4-17](#). The Branches tab displays all commits associated with the branch you are working in. Coalesce also highlights the commit currently associated with your workspace.

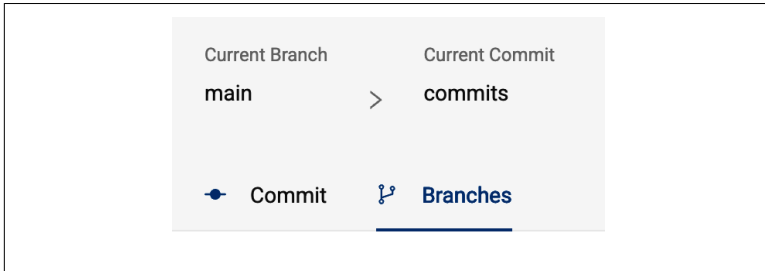


Figure 4-17. The Branches tab, which allows you to work with any branch in your Git repository

To change branches, open the Selected Branch dropdown and choose any branch associated with the Git repository connected to your project, as shown in [Figure 4-18](#). This allows you to pull in branches from other workspaces within the same project.

To check out a branch, select it from the dropdown and click Check Out Latest. This checks out the branch, provided there are no unsaved commits in your current branch. If you need to check out a branch without committing or saving your current work, select the new branch and click Force Checkout. Coalesce will prompt you to confirm this request before it proceeds, as there is a risk of data loss associated with this action.

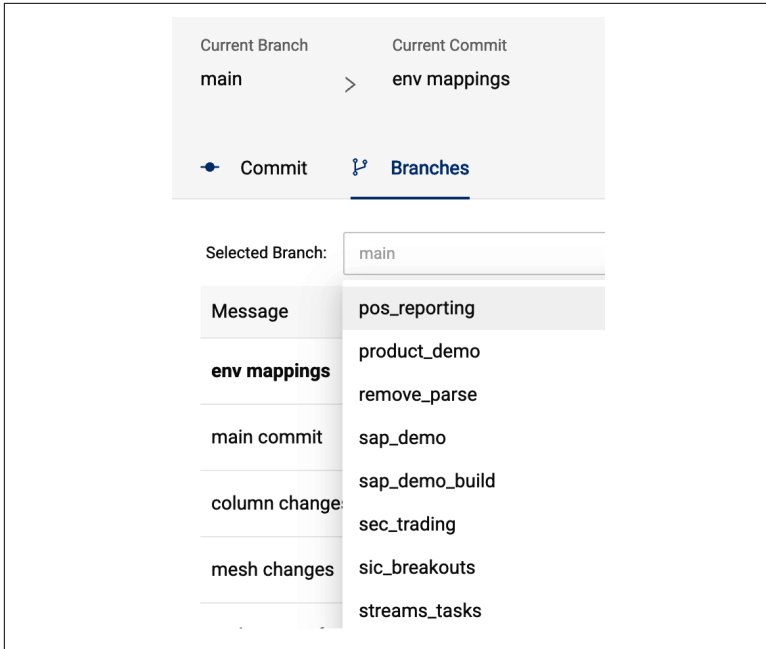


Figure 4-18. The branch selector within Coalesce

You can merge commits from one branch into another directly within Coalesce. To do this, open the Git modal, select Branches, and choose the branch you want to merge your work into. Coalesce will display all commits in the selected branch. Click Merge next to the commit where you want to merge your work. Coalesce will preview the merge before applying it. If everything looks correct, confirm by selecting Merge, and Coalesce will handle the process automatically.

Alternatively, you can use the Merge Latest button to merge only the most recent commit from your current branch.

You can also create a new branch from any commit in any branch within the Git modal. To do this, click New Branch next to the desired commit, as shown in [Figure 4-19](#).

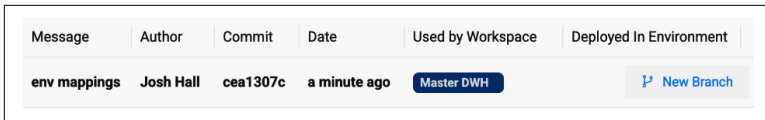


Figure 4-19. You can create a new branch from any commit by selecting the New Branch button

All of the changes within the Git modal in Coalesce will be reflected in the version control system that you have integrated. This gives you complete control over the entire development and versioning of all of your work in the same system.

Macros and Parameters

This section could have been included in the previous chapter on building pipelines, but it fits better here, under managing pipelines, since macros and parameters streamline data development and provide greater control over transformations. Although macros and parameters are related, they serve different purposes, so we will cover them separately. Let's start with macros.

Macros

Macros are functions that take arguments and return a string. You can use these strings in node templates, transformations, and joins to ensure consistency across your project. The returned strings ultimately represent SQL commands, allowing your team to define reusable code in a single location and apply it throughout the project.

For example, to determine if a numerical column contains an odd or even value, you could write a macro as follows:

```
{%- macro even_odd(column) -%}
    CASE WHEN MOD({{ column }}, 2) = 0 THEN 'EVEN' ELSE 'ODD' END
{%- endmacro %}
```

In this macro, the function receives a column as an argument and returns either 'EVEN' or 'ODD'. To use this macro in a transformation, call it using the syntax: `{{ function('"schema"."table"') }}`. In this case, the function is named `even_odd`, so calling the macro would look like this:

```
{{ even_odd('"CUSTOMER"."C_NATIONKEY"') }}
```

Or like this:

```
{{ even_odd('{{SRC}}') }}
```

Macros support everything from simple SQL statements to complex transformations. You can apply them directly to your workspace or include them in any package installed from Coalesce Marketplace.

To add a macro, navigate to the Build settings and select Macros, as shown in [Figure 4-20](#). You will see sections for Workspace Macros and, if installed, Package Macros.

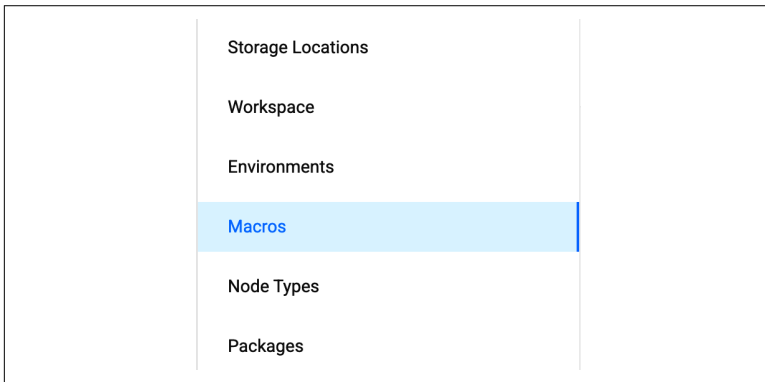


Figure 4-20. Selecting Macros within the Build settings

Sometimes macros need to receive input from a parameter, so let's talk about this next!

Parameters

Parameters act as read-only environment variables. A parameter is defined as a JSON blob set by the user that can be accessed in the metadata during template rendering. Parameter values can be set for a workspace/environment at various points to provide flexibility of template behavior during deploying or refreshing.

There are several reasons to use parameters:

Environment configuration

Parameters allow you to configure your data pipeline for different environments (e.g., development, staging, or production) without modifying the underlying code. For example, you can use parameters to specify different database credentials, file-paths, or other environment-specific settings.

Dynamic data filtering

You can use parameters to dynamically filter or partition data based on certain criteria. For example, you can use a date parameter to only process data for a specific date range or a customer ID parameter to process data for a particular customer.

Reusability and modularity

By changing certain values or logic into parameters, you can make your data pipeline more modular and reusable. This promotes code reuse and maintainability, as you can easily swap out parameter values without modifying the core transformation logic.

Testing and debugging

Parameters can be helpful for testing and debugging purposes. You can use different parameter values to simulate various scenarios or edge cases, making it easier to identify and fix issues in your data pipeline.

Scheduling and automation

When scheduling or automating data pipeline runs, you can use parameters to pass dynamic values or configurations. This allows you to adapt the pipeline behavior without changing the underlying code, enabling greater flexibility and control.

Separation of concerns

By externalizing certain values or configurations as parameters, you can separate the concerns of data transformation logic from the specific values or configurations used in different contexts. This promotes better code organization and maintainability.

Compliance and auditing

In regulated industries or environments with strict data governance requirements, you can use parameters to enforce data access controls, data masking, or other compliance-related configurations.

You can easily add parameters into your workspace or environment. To add a parameter to your workspace, navigate to the Workspace settings and select the Parameters item, as shown in [Figure 4-21](#). You will learn about environments and where to add parameters to them at the end of this chapter ([“Environments” on page 87](#)). When you add parameters to your workspace, they will not be stored within Git.

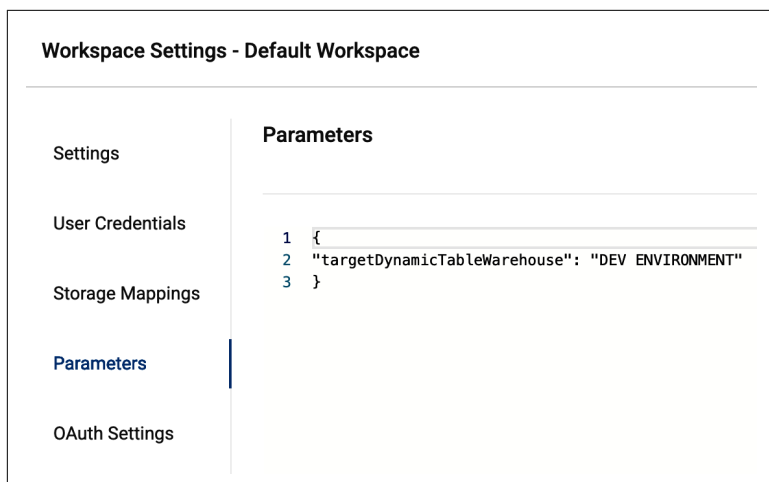


Figure 4-21. The Parameters item within the Workspace settings

Macros and parameters provide a flexible and modular way to organize and manage data and operations throughout your pipelines. Now that you understand how to manage and work with pipelines, the next steps are testing and running them as jobs in an environment. The rest of this chapter will focus on these final steps.

Testing

Coalesce provides built-in testing capabilities to assess data quality. *Data quality tests* help identify issues such as missing values, incorrect data types, outliers, and violations of business rules. Testing within Coalesce ensures that transformed data meets required standards before it is loaded into the target node. Integrating data testing into the pipeline workflow enables early issue detection and ensures consistency across data sources and transformations.

Coalesce handles testing through built-in node-testing capabilities, which support tests at both the column level and the node level, as shown in [Figure 4-22](#).

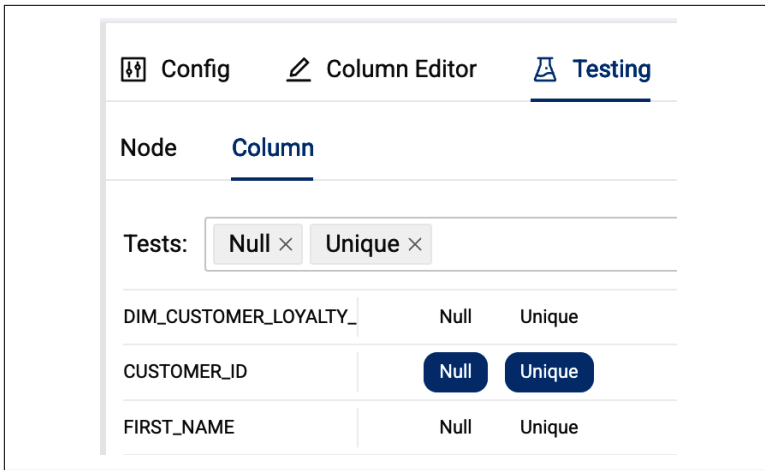


Figure 4-22. Column-level tests available out of the box

You can use these for:

- *Column-level tests:* Test individual columns for uniqueness and null values.
- *Node-level tests:* Validate data quality across an entire node.

To apply tests, select a node and click Testing. Choose the tests to run and specify which columns to include ([Figure 4-22](#)).

You can also write custom node-level tests using SQL queries. If a query returns a result, the test fails. Additionally, you can configure options to determine whether a test failure should halt pipeline execution and whether tests should run before or after the node refresh.

For more comprehensive and extensible testing, you can use the Test Utility Package from Coalesce Marketplace. This package provides granular control over every component of each node and column within your data pipeline. You can learn more about the Test Utility package in the [Coalesce documentation](#).

Jobs

Now that you understand how to manage pipelines, use version control, and test data, the next step is to organize your pipelines into jobs. In Coalesce, you'll use *jobs* to refresh data pipelines. You can

execute these through the CLI or the native Coalesce job scheduler. The CLI is not in the scope of this guide, but you can learn more by viewing the [Coalesce documentation](#).

Each job is identified by a unique Job ID, which includes all nodes associated with that job. You can use a selector query to assign specific nodes from your pipeline to a job. To manage jobs, navigate to the Jobs menu in the left navigation bar. Here, you can view and modify existing jobs and add new jobs. To create a new job, click the plus (+) button and select Create Job. You will then define the job's associated nodes using a selector query, as shown in [Figure 4-23](#).

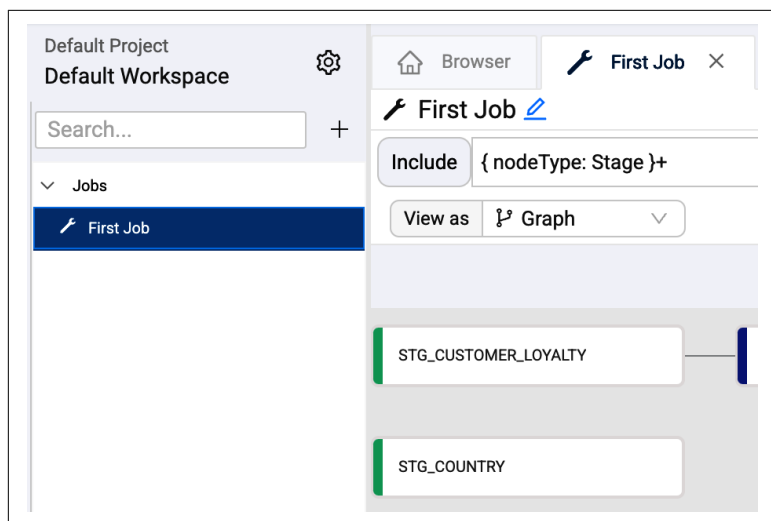


Figure 4-23. Selector syntax for defining a job

Using the job scheduler, you can reference the job name (Job ID) and set up the job to refresh on a schedule. You'll learn more about the Coalesce job scheduler in the next section on environments.

Environments

So far, this chapter has covered how to navigate, version, and manage data pipelines. The final and most critical component of pipeline management is validating pipelines and making them available for consumption. This is where environments come into play.

Environments handle deploying data pipelines to your data platform. They enable efficient deployment strategies, allowing you to

develop in a workspace and then deploy to any environment, such as production or QA. Each environment can have different storage locations, storage mappings, parameters, and other configuration settings.

An environment is configured similarly to a workspace, with its own storage locations and storage mappings and a connection to your data platform. You can view existing environments from the deploy interface in Coalesce. If no environments exist, you can create one from the Build settings of your workspace.

To create a new environment, open Build settings in any workspace and navigate to Environments. View the list of available environments for management and then click Create Environment to open a configuration modal. Provide the connection details for your data platform and define the storage mappings based on the storage locations available in your workspace.

By configuring environments this way, you ensure that data and objects persist within specific locations in your data platform, allowing your team to explore and validate pipelines.

For example, let's assume you have completed development on a new pipeline and need to perform QA testing before deploying to a production environment. You can create a QA environment in Coalesce that deploys pipeline objects to a dedicated database and schema in your data platform for testing.

This approach ensures that your data platform is segmented for different stages of pipeline development, from initial testing to final production release, allowing business users to access validated data with confidence.

Deploying

Deployment is the process of pushing your work to your data platform at the specified location. Once you have created your environments, you can deploy your work at any time to the environment of your choice, as long as the required data sources are included to match those used in the workspace.

To deploy a pipeline to an environment, navigate to the deploy interface and select the blue Deploy button, as shown in [Figure 4-24](#). This opens the deployment modal, where you'll choose a branch from the project you are working in and select a commit to deploy.

For example, if you have just added new nodes to your pipeline and committed the work to your branch, you would select that commit for deployment.

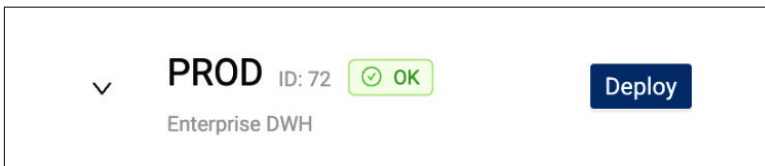


Figure 4-24. Deploying your commits to an environment using the blue Deploy button

Once you select a commit, the next step is to provide any parameters, just as you would when setting parameters in a workspace. Finally, Coalesce will generate a deployment plan, displaying all changes and updates. If there are no errors and the deployment looks correct, confirm the deployment, and Coalesce will push the changes to your data platform.

Refreshing Your Pipeline

Once you have deployed your pipeline, you need a way to refresh the data in those objects. Coalesce allows you to refresh pipelines using either the Coalesce CLI or the native Coalesce scheduler. Again, this guide focuses only on the Coalesce scheduler.

Earlier in this chapter, you learned about jobs ([“Jobs” on page 86](#)). Once you’ve created a job, you can use the Coalesce scheduler to refresh it on a regular cadence, ensuring your pipelines always contain fresh data.

To create a job schedule, navigate to the deploy interface and select the ellipsis next to the blue Deploy button, as shown in [Figure 4-25](#). From here, choose View Scheduled Jobs to see all available jobs. You can also create a new job schedule by selecting Create Job Schedule in the upper right corner.

NOTE

You must be an environment admin to create and schedule a job. We’ll cover role-based access control (RBAC) in the next chapter.

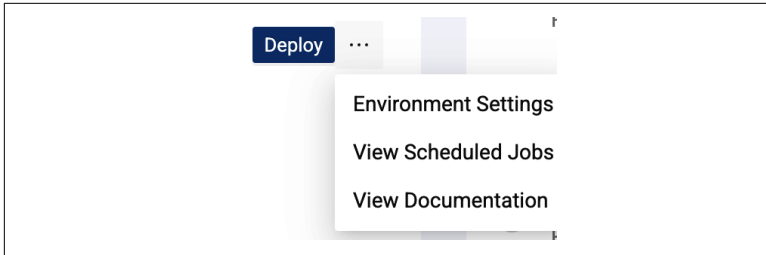


Figure 4-25. Scheduling a job by clicking the three dots next to the Deploy button

To complete the job schedule setup, fill in the required parameters, including project, environment, and the job to run. Finally, specify the schedule for the job execution. Once configured, the job schedule will automatically refresh all nodes associated with the job at the specified intervals.

By putting your job on a schedule, you can have your nodes refresh in your specified environment at the cadence of your choice. Instead of having to perform development time actions, like clicking the Create and Run buttons, you can have this managed for you by deploying your work and then refreshing (running) your pipeline using the scheduler.

The Makings of a Pro

In this chapter and the previous one, you have learned how to build, manage, and refresh pipelines in Coalesce. Specifically, this chapter covered how to navigate data pipelines, create subgraphs, use version control, leverage column-level lineage, and configure environments. With these skills, you are fully equipped to build and manage any pipeline in Coalesce. You are now well on your way to becoming a Coalesce pro and fully self-sufficient in end-to-end pipeline development.

In the next chapter, you will explore security and governance in Coalesce. This will give you the ability to not only build and manage pipelines but also oversee and secure your entire Coalesce instance as an administrator. See you there!

Coalesce Security and Data Governance

Throughout this guide, you've laid the foundation for your data projects. You've used Coalesce's building blocks to construct your house, and in the previous chapter, you organized everything inside your home. Now it's time to protect what you've built. In this chapter, you'll focus on security and data governance in Coalesce—locking the doors, securing the windows, and keeping everything safe.

By the end of this chapter, you'll know how Coalesce supports multi-factor authentication (MFA) and single sign-on (SSO). You'll learn how to configure OAuth for secure authentication. You'll also understand how role-based access control (RBAC) works and how to set it up for your organization. Finally, you'll learn how to get a clear view of everything you've built and see how it connects across your entire data stack. Let's dig in!

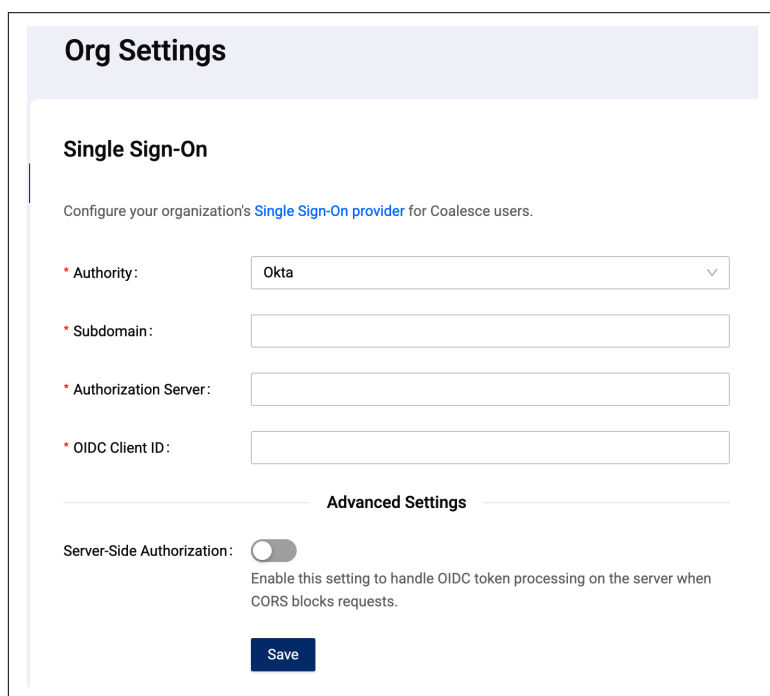
Account Access

Let's begin the conversation about security by talking about accessing your Coalesce instance. Coalesce provides different methods to authenticate, or log in to, your Coalesce account. As of the time of this writing, those methods are to enter a username with a password and to use a variety of supported single sign-on providers, such as Okta. Let's go over single sign-on a bit more.

Single Sign-On

Single sign-on allows users to authenticate into Coalesce using credentials managed by an external identity provider, streamlining access and improving security. Coalesce supports multi-organization environments, so SSO users can log in to multiple organizations with the same credentials. SSO users are provisioned “just in time,” meaning their accounts are created automatically when they first log in through your identity provider and no manual setup is required beforehand.

To configure SSO in Coalesce, go to the user menu in the upper right corner and select Organization settings—just like you did in [Chapter 1](#). From there, choose Single Sign-On and then select your SSO authority, as shown in [Figure 5-1](#). Follow the instructions in the Coalesce [documentation](#) specific to your SSO provider to complete the setup.



The screenshot shows the 'Org Settings' interface in Coalesce. The 'Single Sign-On' section is active, displaying instructions to configure the organization's SSO provider. Below the instructions are four input fields: 'Authority' (a dropdown menu with 'Okta' selected), 'Subdomain', 'Authorization Server', and 'OIDC Client ID'. An 'Advanced Settings' section follows, featuring a 'Server-Side Authorization' toggle switch (currently off) and a descriptive text: 'Enable this setting to handle OIDC token processing on the server when CORS blocks requests.' A blue 'Save' button is located at the bottom of the form.

Figure 5-1. Single sign-on settings within Coalesce

While SSO is a convenient way to manage access, it's not the only option. If you prefer to use a username and password to log in to Coalesce, you can add another layer of security by enabling MFA.

Multi-Factor Authentication

Multi-factor authentication adds an extra layer of security to your Coalesce account by requiring multiple factors: something you know (your password) and something you have (a time-based code from an authenticator app). Even if your password is compromised, MFA helps prevent unauthorized access by ensuring that only someone with access to your second factor can log in.

To enable MFA in Coalesce, first navigate to User settings by clicking the user menu in the upper right corner. Then select Account and Security. Before turning on MFA, make sure your email address is verified. If it's not, resend the verification email and follow the prompts. Coalesce suggests using an authenticator app as your MFA method. Once you've selected MFA, scan the QR code from your mobile device and you'll see a six-digit code appear for Coalesce in your authenticator app.

Keep in mind that, once MFA is enabled, it can't be turned off by the user. Only an administrator can disable it. MFA is managed on a per-user basis and can't be enforced or disabled across the entire Coalesce organization. Each user must enable MFA individually.

So far, we've covered how users authenticate and access Coalesce. But how do you control *what* they can access once they're inside? That's where RBAC comes in, and that's what we'll cover next.

Role-Based Access Control

Role-based access control in Coalesce gives organizations fine-grained control over who can access data assets and what actions they can perform within the platform. RBAC works by assigning users to roles. Each role carries a defined set of permissions that determine access to environments, objects, and specific capabilities, such as deploying jobs or editing nodes. Instead of assigning permissions to individual users, RBAC simplifies user management by allowing administrators to create and manage roles, making it easier to scale access control as teams grow.

At its core, RBAC helps organizations enforce data governance and security policies. It ensures that users only have access to the areas of Coalesce they need to perform their tasks. For example, a data engineer working in development might need full access to create and edit nodes, while an analyst working in production may only need read access to published tables. RBAC helps you enforce these boundaries in a clear and auditable way.

Coalesce RBAC is applied at the organization level and then again at the project and environment levels. This allows for granular control over exactly who should be accessing what projects and data, as shown in [Figure 5-2](#).

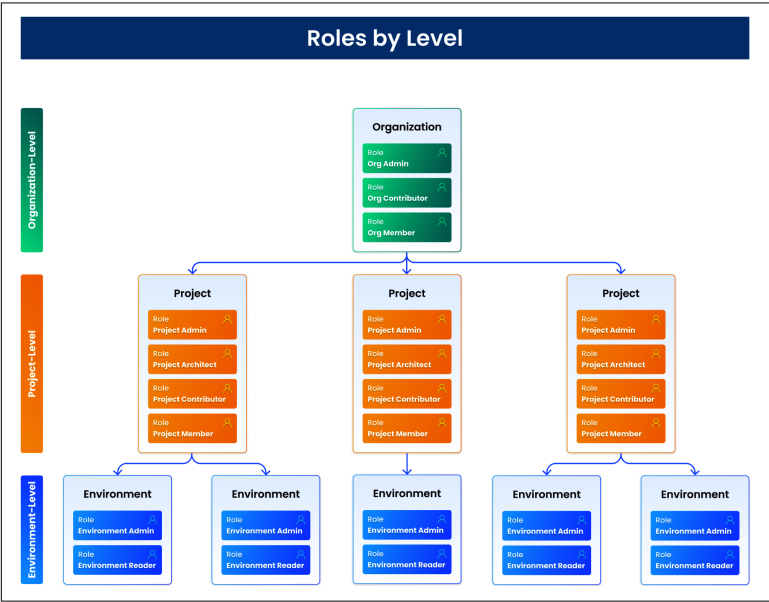


Figure 5-2. The hierarchy of role levels within Coalesce

At the organization level, assign users an org-level role: admin, contributor, or member. After that, you can assign them to specific projects and environments with roles like admin, contributor, architect, or member. Finally, you have the option to add users to environments as either an environment admin or an environment reader. You can see a detailed breakdown of each of these roles in [Tables 5-1 through 5-3](#).

Table 5-1. Organization-level roles

Role	Permissions summary	Recommended for . . .
Organization administrator	The creator of the Coalesce app is automatically assigned the organization administrator role. Only organization administrators can add other users, including other organization administrators. Admins have full access to all functionality in Coalesce.	Individuals who need full administrative control
Organization contributor	Contributors can't add new users to the organization. They have access to read documentation and create API tokens, user settings, and Git account information. They will be able to set up a project, configure Git, add members to projects, and oversee work. They have access only to projects they have created.	Managers who decide how each person will contribute to a project
Organization member	This is the default role. Members can edit Git account information, create API tokens, and read documentation.	Default role

Table 5-2. Project-level roles

Role	Permissions summary	Recommended for . . .
Project administrator	This role can manage projects but cannot create them. The ability to create projects is granted to the organization administrator or contributor role. Project administrators have access to projects, deployments, and environments.	Team managers or senior team members who will manage projects
Project architect	This role can manage certain project information, build nodes, and generate API tokens.	Senior data architects who will be building node types, setting storage locations, and creating macros
Project contributor	A project contributor has read-only access to certain project settings and read access to projects, deployments, and environments. This role cannot edit or create custom nodes or macros.	Team members who need to access project information without making changes
Project member	Project members are added to the environment.	A data engineer or data platform engineer—someone who will not be actively involved in creating or maintaining data pipelines but will need access to the environment level

Table 5-3. Environment-level roles

Role	Permissions summary	Recommended for . . .
Environment admin	This role manages environment settings, reads project documentation, and deploys through the API, CLI, or Coalesce app.	A data platform engineer or an operations engineer who will approve deployments and schedule jobs
Environment reader	This role has access only to the documentation for the environment they are added to. They also have access to certain API functions to get deployment information.	A business analyst or data analyst

To manage roles, start by opening the user menu in the upper right corner of the Coalesce interface. From there, select Organization settings and then click Users—you’ll see a list of existing users. You can create a new user and assign the organization role at the time of creation. Alternatively, you can select the Actions button to edit the user, updating their organization role, as shown in [Figure 5-3](#).

Add User

* First Name: John

* Last Name: Doe

* Email: john.doe@email.com

Set User Password ⓘ : ☐

* Role:

Org Admin
Full admin rights, can edit Org Settings

Org Contributor
Create new Projects, view and edit Projects

Org Member
View and edit Projects

[Project Roles in the Coalesce Docs](#) ⓘ

Josh Hall

Josh Test

Figure 5-3. Assigning an organization role to a user in Coalesce

After a user has an organization-level role assigned, you can assign them to projects and environments. On the Projects page, select any project within your Coalesce organization and click the Project Settings button. You'll see the Members item within the modal, and you can select this to view the users in your Coalesce organization who have been added to the project.

To add a user to the project, select Add Member, choose their name, and then provision them with the project-level access you want them to have, as shown in [Figure 5-4](#). Once you've added a member, you can easily remove them by clicking the Remove button next to their name. Managing access to projects is that easy.

Add Project Member

*** Name**

John Doe

You can only add members that are already part of your Organization at Coalesce.

*** Role**

- Project Admin**
Edit Project, Workspaces, Macros and UDNs, and Nodes
- Project Architect**
Edit Workspaces, Macros and UDNs, and Nodes
- Project Contributor**
Edit Workspaces and Nodes
- Project Member**
No Project and Workspace permissions, view permissions of environments they are explicitly given access to.

[Project Roles in the Coalesce Docs](#)

Figure 5-4. Assigning a user a project-level role

You can follow a similar process to provision users to environments. Select the deploy interface and then the ellipsis next to the environment you want to provision access to. Select Environment settings and, similar to project access, you will see the Members item within the modal. Within Members, select Add Member to add a user who has been provisioned with an organization-level role, as shown in [Figure 5-5](#).

Add Environment Member

* Name

John Doe

Only project members can be added to an environment. Learn more in the [documentation](#).

* Role

Environment Admin
Edit, deploy and refresh Environments. View Run Results and Documentation

Environment Reader
View Run Results and Documentation

[Environment Roles in the Coalesce Docs](#)

Figure 5-5. Adding a user to an environment in Coalesce

RBAC is a key part of maintaining a secure and well-governed Coalesce environment. It ensures that users can interact only with the data and environments that are relevant to their role in the data pipeline lifecycle. But in some cases, it's not just your users that you want to provision to ensure your data is secure—it's the technology itself. In the next section, we'll talk about how Coalesce executes your pipelines on your data platform in a secure way.

Coalesce SQL Execution

Ensuring the technologies that you are using within your organization are compliant with your security requirements is critical. As a solution that transforms data, Coalesce has access to metadata within your data platform. However, Coalesce does not store data from your data platform. Ever. Coalesce strictly executes the SQL that is generated based on the metadata from your objects and the SQL transformations you have supplied.

To support SQL execution from within a browser environment, Coalesce uses a direct proxy that forwards SQL text on behalf of user actions submitted through the Coalesce UI to the customer's data platform instance. All data is encrypted via HTTPS/TLS, and data results from SQL executions are never stored by Coalesce.

When you submit a SQL execution, the browser sends the SQL text to the Coalesce backend. After authenticating the user, Coalesce retrieves credentials from the industry-standard credential manager and uses them to connect to your data platform. It then sends the results back to the browser over HTTPS and TLS.

When you request to fetch data, Coalesce follows the same SQL execution path and returns the first hundred rows to the browser for display. Coalesce never stores the data.

As you have seen, Coalesce never stores the data it processes. It only passes encrypted SQL and, when needed, returns a limited result set, one hundred rows, to the browser.

With your users now provisioned correctly and your security in place, it is time to focus on sharing the results of your data projects. How do you expose that value to your stakeholders? That is exactly what the rest of this chapter will explore.

Documentation

Though it is an essential component of building a sustainable data architecture, documentation is sometimes put aside to focus on meeting the deadlines for development work. Fortunately, Coalesce automatically produces and updates documentation as developers work.

At the core of Coalesce’s documentation process is its ability to capture everything about each node you have created within any project, workspace, and environment. Information such as the data-base and schema of the object, columns, data types, lineages, and even the DDL and DML of each node is captured automatically.

However, Coalesce users can elevate their documentation further through the use of column- and node-level documentation.

Column-level documentation lets users add descriptive comments and context for every column within a node. You can enter this information in the Description column within the mapping grid of a node, as shown in **Figure 5-6**. This field accepts plain text, which you can either enter manually or generate automatically using the Coalesce AI description tool.

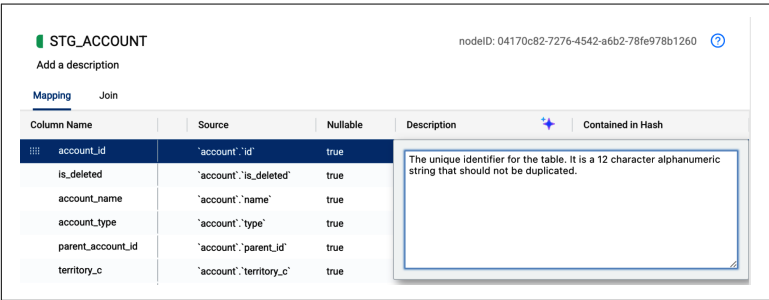


Figure 5-6. Providing a column-level description within Coalesce

The same applies to node-level documentation. Business users often find clear descriptions of objects helpful, including what the objects contain and how they are intended to be used. In Coalesce, users can document a node by selecting the “Add a description” text box located beneath the node name, as shown in [Figure 5-7](#). Just like with column-level descriptions, users can enter this information manually or use the Coalesce AI tool to generate a description automatically.

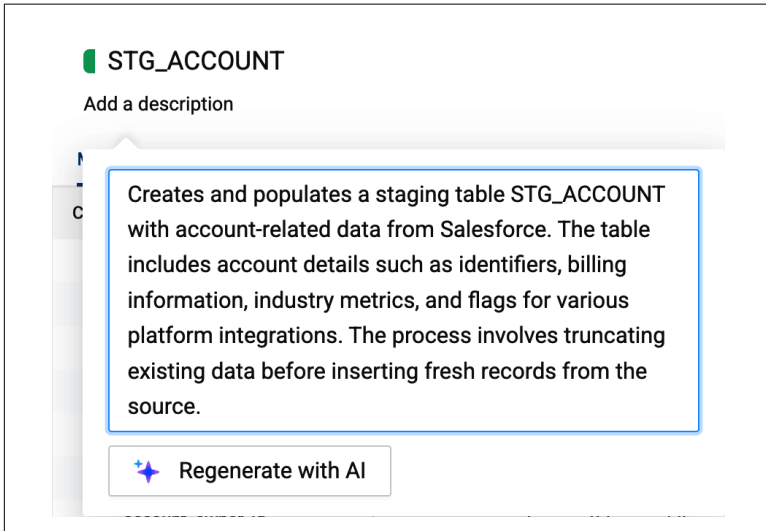


Figure 5-7. Node-level description within Coalesce

Coalesce captures all of this documentation in real time and displays it in the docs interface. When you open the docs interface, you can view all projects, workspaces, and environments across your Coalesce organization. Selecting any workspace or environment will show you all of the nodes within that location, as shown in [Figure 5-8](#), along with the metadata associated with each node.

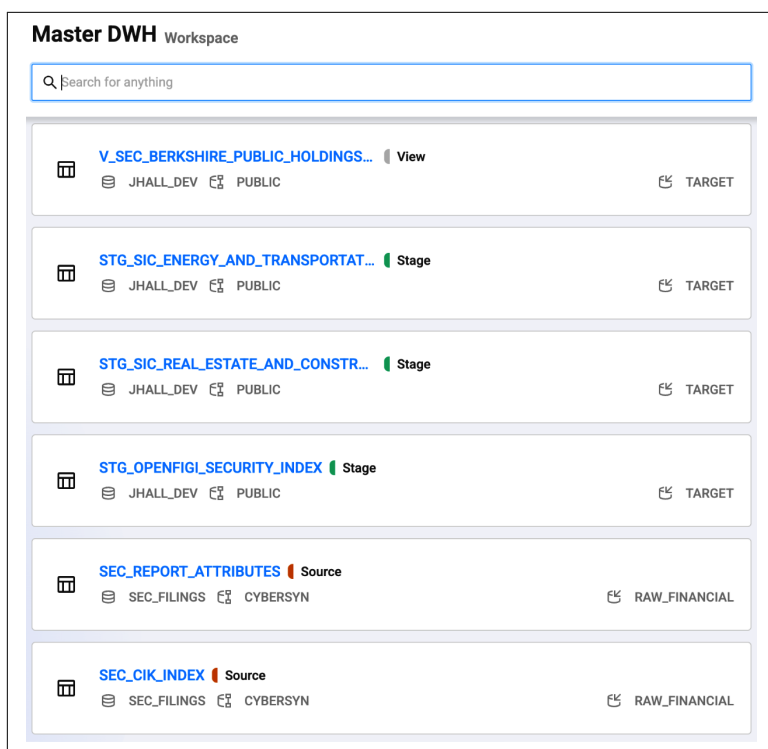


Figure 5-8. Line items representing different nodes that have been documented automatically in Coalesce

By selecting a node, you can drill down to a more detailed view of its contents, as shown in [Figure 5-9](#). The docs interface also includes filtering options that make it easy to find specific nodes or pieces of documentation. This gives users the context they need to answer questions, troubleshoot issues, and understand how data flows through the project.

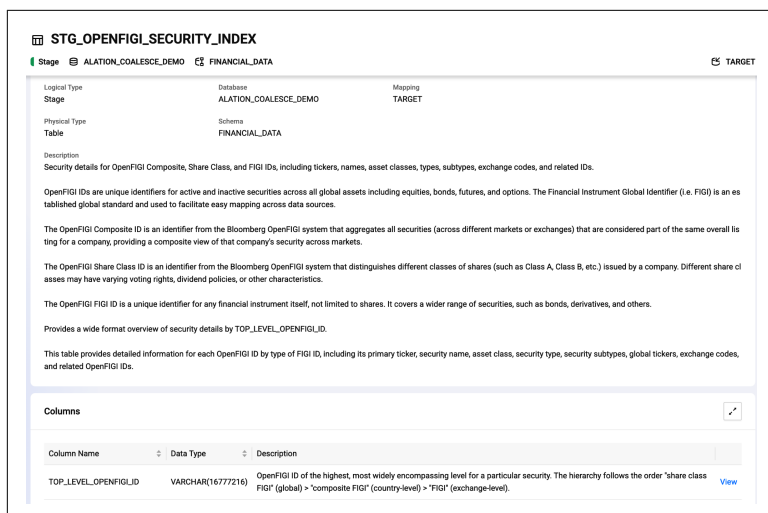


Figure 5-9. Documentation within a node inside the docs interface

Coalesce lets users manage their organization, control access, and view activity across all projects.

Peace of Mind

Coalesce includes intentional security measures at every step to protect your organization, your data, and your users. In this chapter, you learned the different ways users can authenticate into Coalesce. You also saw how to control what users can access once inside and how Coalesce ensures that your data is never stored, all while making your pipelines easy to understand through clear, structured documentation.

At this point, you have mastered the core components of Coalesce. From setting up your workspace to building pipelines, managing projects, and enforcing security and governance, you now have a complete view of what Coalesce can do. In the next chapter, you will put that knowledge to work as you explore some of the real-world use cases that Coalesce helps organizations solve.

Modeling Patterns and Use Cases in Coalesce

By now, you have set up your Coalesce account and built scalable, production-ready pipelines. You have handled the core tasks: transforming, versioning, testing, and deploying data transformations. The next question is, what are you building for?

The answer depends on your organization's goals, but most teams follow a few common architectural patterns. There is no universal blueprint, but there are proven approaches that guide how teams model and operate at scale.

Throughout this book, we have used a house-building analogy to ground each phase of development. You laid the foundation, framed the structure, and secured your belongings. Now it is time to decide how to live in the space you created. How should the rooms be arranged? What temperature should the thermostat be set to? What rooms should the kids sleep in? Anyone who moved into this house would have to decide these things, but everyone might do it a little differently. That's exactly what we'll be focusing on in this chapter.

This chapter shifts the focus from how you build to why, and what comes next. It explores common patterns in modern data environments such as migrations, dimensional modeling, and data vault architecture. Each one addresses a different type of problem and requires a specific approach to modeling, testing, and maintaining pipelines over time. By the end of the chapter, you'll have a better understanding of the types of problems you can solve in Coalesce

and how to take advantage of powerful off-the-shelf functionality to solve your toughest problems.

Migration: From Legacy to Modern

Migrations usually start with a mix of dread and necessity. Moving SQL logic from one platform to another often means rewriting, refactoring, and revalidating everything. It's slow. It's tedious. And it's easy to miss things. But there are ways to break it down. Think of migration as being like moving to a new house. You have several options:

- You pack your own boxes, use your own truck, and move everything yourself.
- You hire movers to pack, drive, and unpack for you.
- You throw out or sell everything and start fresh.

All of these options are valid. Your approach depends on what you're moving, how urgent the move is, and how much help you have. In data terms, that means deciding whether to rebuild, lift-and-shift, or rearchitect.

Building New

There is a unique kind of momentum that comes with building from scratch. No legacy models to untangle. No outdated logic to reverse engineer. Just a clean starting point and the freedom to design something that actually fits your data and your team.

Coalesce makes it easy to start fresh. With nodes, teams can apply consistent modeling patterns across your architecture without rewriting the same logic over and over. Because nodes are customizable, they reflect your team's standards, not a fixed set of rules. Utilizing both nodes and Coalesce Marketplace makes it straightforward to define repeatable structures that still allow for edge cases.

For example, say I'm developing on Snowflake and want to take advantage of Iceberg and Cortex functionality. Instead of writing those templates myself or, worse, having to figure out how to get that functionality working, I can start building my pipelines using packages from Coalesce Marketplace that have already defined that functionality.

Because Coalesce separates node-type configuration from the SQL logic you write to transform your data, the process of building out data models stays organized. You are not jumping between files or scanning through long SQL scripts to find what changed. Instead, changes are visual, versioned, and easy to understand.

One of the clearest advantages of starting from scratch in Coalesce is that you avoid inheriting technical debt. Every new node benefits from built-in lineage, clear naming conventions, and automatic documentation. The structure you define on day one scales with the project, rather than breaking down as complexity increases.

This is not about moving fast at the expense of long-term stability. Coalesce makes it possible to move quickly because the system is designed to prevent chaos later. When you build new, you are creating pipelines that are easier to maintain, easier to test, and easier for others to understand.

If you have the opportunity to start clean, Coalesce gives you the structure and flexibility to get it right the first time without slowing down.

Migrating As Is (Lift-and-Shift)

A lift-and-shift is often the first step in a larger transition. Instead of redesigning everything from the start, you focus on getting your existing pipelines running in the new environment. It is a pragmatic move, especially when time, complexity, or risk makes a full rebuild unrealistic.

In Coalesce, a lift-and-shift means re-creating existing logic in a cleaner, more structured way. You bring in the SQL that already works, but you do it with better guardrails. Clearly defined nodes structure your SQL, streamlining business logic. And your pipelines reflect how the data actually flows, not how the flow has evolved over time.

You are not inventing something new here. You are giving your current data process a stronger foundation. Many teams use this phase to wrap legacy transformations in Coalesce nodes, apply consistent patterns, and start building toward more modular design, even if the core logic remains the same.

This is also the time to clean up the small things that cause friction. Confusing joins, unexplained filters, cryptic column names . . . when these show up, Coalesce gives you the structure and visibility to fix them without losing track of the original intent.

Migrating as is does not mean standing still. By moving existing pipelines into Coalesce with clear structure and built-in governance, you create a stable baseline for future changes and avoid carrying over the same problems into the next phase.

Rearchitecting for the Future

Rearchitecting often follows a lift-and-shift. Once everything is running in the new environment, teams start to notice what no longer works. The logic that made sense five years ago may not hold up today. Business requirements change. Data volumes increase. What used to be manageable becomes brittle, opaque, or too difficult to extend.

In Coalesce, rearchitecting does not have to mean starting over. You can keep what works and replace what doesn't. Most teams begin by creating clearer boundaries—separating the raw, staging, and refined layers, for example, or breaking monolithic SQL into smaller, purpose-built nodes. Business logic is fully tested by the team. Naming conventions get a second look. So does the overall shape of the pipeline.

The visual interface in Coalesce makes this easier to manage. You can trace dependencies across your pipeline, spot unnecessary complexity, and understand where changes will have downstream impact. It gives you the context you need to plan a rearchitecture without guesswork.

This kind of work is usually incremental. Few teams rewrite their pipelines all at once. Instead, they modernize one piece at a time, refactoring logic, reorganizing layers, and cleaning up joins and filters as they go. Coalesce supports this approach by making changes transparent, standardized, and reversible. You can experiment with confidence, knowing every change is tracked and easy to roll back if needed.

Rearchitecting is about getting to a place where your pipeline is easier to understand, easier to maintain, and better aligned with how your organization actually uses data. Coalesce helps you do that

without throwing everything away. As you build in Coalesce, modeling your data for sustainable and maintainable pipelines suddenly becomes important.

Modeling

There is no single way to model data. The right approach depends on your goals, your constraints, and how your team plans to use the data. In the sections that follow, we will walk through several important modeling paradigms: dimensional modeling, data vault, and operationalizing AI and ML pipeline development and maintenance with nodes and node packages. Each one reflects a different way of thinking about structure, whether you are optimizing for performance, traceability, or predictive insight. Coalesce supports all of these approaches by giving you the tools to apply consistent logic, enforce standards, and adapt your pipeline architecture as your needs evolve.

Dimensional Modeling: The Foundation of Analytics

Dimensional modeling has endured because it provides clarity. It structures data in a way that analysts, engineers, and business users can all understand. At its core, dimensional modeling separates data into two types of tables: facts and dimensions.

Fact tables record measurable events such as orders, page views, revenue, or inventory levels. They are typically long, narrow tables with foreign keys that point to dimension tables. *Dimension tables* provide context. They describe the who, what, where, and when behind each fact, such as customers, products, locations, dates, and so on. The result is a model that is predictable and easy to navigate.

There are two major variations of dimensional modeling: star and snowflake. A *star schema* keeps dimensions directly connected to facts. A *snowflake schema* organizes dimensions further into related lookup tables, often to reflect hierarchies or reduce redundancy. Either approach gives downstream consumers a consistent structure and allows business intelligence tools and SQL engines to optimize queries more effectively.

Coalesce supports dimensional modeling not just as a concept but as a core design pattern. You can explicitly define Fact and Dimension nodes and apply consistent standards across both. This is not

just a label: Fact and Dimension nodes behave in structured ways. They follow naming conventions; align to logical layers; contain optimized, best-practice SQL; and reflect real relationships in your data. The Coalesce interface gives you clear visibility into how these pieces connect.

Coalesce also helps to enforce modeling best practices without relying on institutional knowledge. The templates for Fact and Dimension nodes reflect your team's standards, whether that means how you generate business keys, apply audit columns, or implement slowly changing dimension logic. Once the template is defined, it becomes part of how your team works, not a checklist you have to remember.

As your data model evolves, this structure is what keeps it from becoming fragile. The strength of a dimensional model comes from its consistency. Coalesce gives you the tools to maintain that consistency while still allowing flexibility when needed. You can make changes with confidence, understand the impact of those changes, and keep everything documented automatically.

So while dimensional modeling is a well-established practice, Coalesce brings it into a modern workflow. With repeatable patterns, clear lineage, automated documentation, and strong version control, you can build and maintain dimensional models that scale with your organization and remain easy to work with over time.

Data Vault: Designed for Flexibility and Traceability

Dimensional models work well when the scope is known and the data is clean. But not every environment fits that mold. Sometimes you are working with dozens of source systems, each evolving on its own schedule. Sometimes the requirements include strict auditability, full historical tracking, or regulatory compliance. In those cases, dimensional modeling alone may fall short. This is where the data vault paradigm becomes essential.

Data vault modeling is designed for flexibility, scale, and traceability. It separates structural elements from descriptive ones, making it easier to adapt as systems and business logic evolve. There are three main components of a data vault system: hubs, links, and satellites. *Hubs* store core business entities using unique business keys. *Links* define the relationships between those entities. *Satellites* capture descriptive attributes and track historical changes over time.

This separation allows you to ingest and structure raw data in a way that does not require up-front assumptions about how the data will be used. You are not forced to model for analytics on day one. Instead, you build a persistent layer of fact that reflects the raw truth of what happened, when it happened, and how it changed.

This is where data vault shines. It is built for change. When you add a new source system or onboard a new satellite, you do not need to redesign your existing structure. You simply extend the model. The core business keys remain stable; the lineage remains intact; and every change is versioned, timestamped, and traceable.

Data vault is not optimized out of the box for reporting speed. That is not its goal. Its strength lies in preserving the integrity and history of your data. It creates a single, trusted layer that downstream models, dimensional or otherwise, can pull from with confidence. In regulated industries, or any environment where data lineage and historical accuracy matter, this becomes critical.

Coalesce enhances that traceability by using lineage that is not a static diagram but contains a fully interactive map of where each field came from and how it was transformed. Combined with automatic documentation, version control, and Git integration, it becomes easy to reconstruct the state of your data model at any point in time.

The modular structure of the vault also makes it easier to manage complexity as your architecture grows. You can incrementally build and refactor hubs, links, and satellites without disrupting downstream models. Coalesce supports this kind of modular thinking, giving you the tools to build a flexible, extensible data backbone without losing control.

So when the priority is not just reporting but long-term data stewardship, Coalesce and data vault are a natural match. They give you the structure, visibility, and discipline to scale confidently, even when everything else is changing, especially in a world where machine learning and artificial intelligence have gone mainstream.

Wrapping Up

Architectural patterns are not about following trends. They are about solving the real problems that arise when data volume grows, systems multiply, and business needs evolve. Whether you are starting fresh, reworking legacy logic, modeling for analytics, preparing AI features, or integrating complex external sources, the patterns in this chapter give you a framework for thinking clearly about structure.

Coalesce is built for exactly this kind of work. It gives you a way to standardize without getting stuck. Templates let you codify best practices without locking you into rigid abstractions. Visual lineage helps you understand what is happening and why. And the ability to wrap even advanced functionality, like AI functionality, into nodes means you are never stitching things together on the side.

The best pipelines are not just technically correct. They are understandable. They are maintainable. And most importantly, they are adaptable. That is what sets long-lived architectures apart: they support change without falling apart.

You learned in this chapter how to think about migrations within Coalesce. You also learned about some of the common patterns that Coalesce seamlessly supports. These patterns are just one way to think about what comes next. Whether you are scaling up, branching out, or just trying to keep things clean as complexity grows, Coalesce gives you the structure to stay organized and the flexibility to move fast.

If you have made it this far, you are likely already solving meaningful problems with data. You've built your data products with confidence. In the next chapter, you'll see how your team uses them, learn how to govern them more effectively, and use that insight to iterate and improve your pipelines.

Wrapping It All Up with Coalesce Catalog

At this point, you have everything you need to own your data transformations—whether you’re a team of one or managing a hundred engineers. In [Chapter 5](#), you locked down governance and security, and in [Chapter 6](#), you reviewed common patterns for data migration and modeling. Now it’s time to take the final step: making your data truly discoverable and usable across your business. That’s where Coalesce Catalog comes in.

This chapter is all about helping you take action, not just on the data you’ve transformed but on the rest of the systems in your stack. Coalesce Catalog is your window into the entire data ecosystem. Think of it like the index of a book . It tells you what’s inside, where to find it, and how everything fits together.

Or, to bring back our house analogy, this is where you stop building and start living. Coalesce Catalog shows you what’s in the fridge, where the remote’s hiding, and how clean the kids *actually* left their rooms. You can even ask it questions and it won’t talk back.

By the end of this chapter, you’ll know how to use Coalesce Catalog to unlock the full value of your transformed data, support your stakeholders, and close the loop between creation and consumption. Let’s get into it.

Getting Set Up with Coalesce Catalog

Before you can explore metadata, digest documentation, and start understanding lineage, you'll need to connect Coalesce Catalog to your data sources. Setup is fast and straightforward. In this section, we'll walk through how to access your workspace and connect to your data stack to start enriching your Catalog from day one.

Once that foundation is in place, you'll be ready to navigate the Catalog, support your team, and empower business users to answer their own questions anytime they want.

Accessing Your Workspace

Coalesce Catalog is a cloud-based application, just like the Coalesce transformation product you've learned about throughout this book. You'll either receive an invite link via an email or be prompted to sign in using your work email or SSO.

During onboarding, Coalesce will typically help configure key settings for your organization—like enabling approved email domains or setting up SSO. Once that's in place, users can self-serve and create accounts with minimal friction.

User roles are intentionally simple, and most users can explore and document your data out of the box. Admins have extra permissions for managing integrations and settings, but there's no complicated access model to learn.

Connecting to Your Data Stack

Once you're in, the next step is to connect your primary data source. Coalesce Catalog supports a wide range of databases and warehouses, but the setup process is consistent no matter which one you use.

Start by creating a dedicated read-only account in your database. This account only needs access to metadata, like schemas, tables, columns, and query history. It should be restricted from accessing any actual data. It's highly advised to follow along with the onboarding steps of the tooling you are integrating with Coalesce Catalog, as each step will have dedicated support for the system you are integrating.

Once the account is ready, open the Integrations section in Coalesce Catalog, as shown in [Figure 7-1](#). Choose a database to connect with and enter your connection details. These typically include a host, credentials, and any necessary configuration values. Credentials are encrypted, and only metadata—not data—is ever ingested.

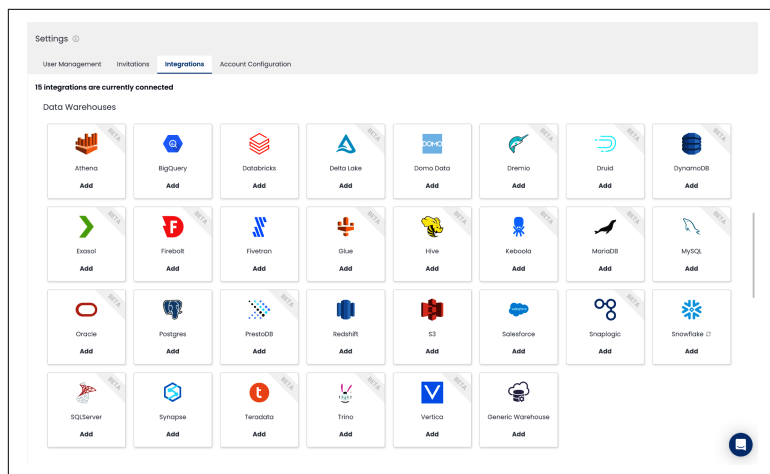


Figure 7-1. The Integrations tab in Coalesce Catalog, where you can integrate with over sixty modern data stack systems

After saving the connection, Catalog runs an initial sync. It scans your database for structure, databases, schemas, tables, and columns and collects query history to begin building usage insights and lineage. If your environment is large, the first sync may take a bit, but you'll be notified once it's complete.

From then on, the catalog stays up to date automatically. Changes like new tables, renamed columns, or updated usage patterns are picked up on the next scheduled sync. If needed, you can fine-tune which parts of your environment are scanned by applying filters to include or exclude specific schemas or tables.

You can follow the same process to connect additional tools—whether that's another database, your BI platform, or something else in your data stack. Just choose the integration, provide the credentials, and let Coalesce Catalog handle the rest.

Metadata Management and Data Lineage

Coalesce Catalog is built to simplify metadata management and illuminate how data flows across your environment. In this section, you’ll learn how metadata is ingested, how to document and enrich your assets, and how to use the lineage graph to bring clarity and confidence to your data.

Documenting Data Assets

Coalesce Catalog makes it easy to enhance your metadata so it’s not just accurate but genuinely useful. Each table, dashboard, and column can include rich-text documentation, as shown in [Figure 7-2](#).

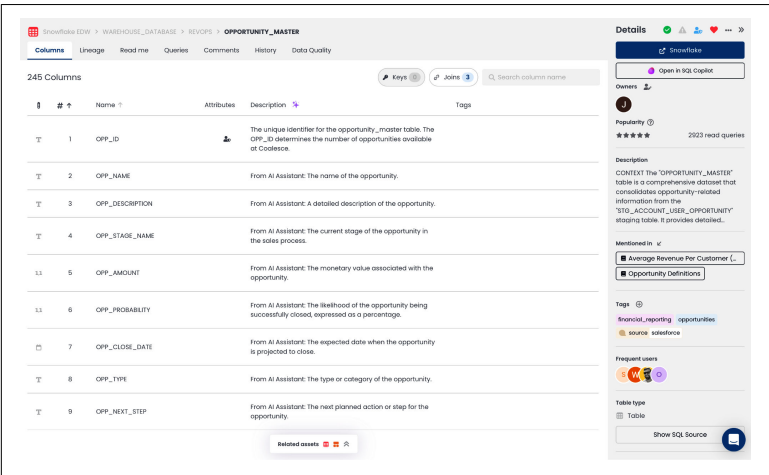


Figure 7-2. Using rich text to document columns and provide context to the table

When documenting common fields like `id` or `created_at`, Coalesce Catalog suggests existing definitions from similar columns elsewhere in your catalog. You can choose to reuse or propagate those descriptions to other same-named columns to maintain consistency. If you need to document many fields at once, the metadata editor provides a spreadsheet-style view to streamline updates.

AI support can help speed up your documentation process. The “describe with AI” feature drafts suggestions based on context and naming conventions. You can review, edit, or regenerate as needed; the AI functionality is designed to help, not replace, your judgment.

Within each asset, you can use tags to categorize assets by aspects such as business domain, sensitivity, or quality status. Then you can assign owners to each asset to ensure accountability. Alternatively, you can take advantage of the automatic surfacing of frequent users so you can identify subject matter experts, even if they aren't officially assigned.

To help teams track progress, Coalesce Catalog also calculates a *documentation completeness score*. This score is based on criteria like whether an asset has an owner, a table description, and column-level documentation. While all of this documentation is helpful for gaining context, one of the best ways to understand how data assets came to exist is by understanding their lineage.

Understanding Data Lineage

Lineage is one of the most powerful features in Coalesce Catalog. It shows how data moves from source to transformation to consumption—helping you debug issues, validate assumptions, and avoid downstream breakage.

You can explore lineage in two main ways: the list view and visual graph. The *list view* gives you a summary of upstream and downstream relationships. The *visual graph* brings everything to life, as shown in [Figure 7-3](#). It starts focused on just the selected asset and lets you incrementally expand upstream or downstream by clicking the plus (+) or minus (–) buttons. This makes it easy to trace dependencies or investigate impact without getting overwhelmed. Different asset types are visually distinguished, and interactive elements help you move through the lineage step-by-step.

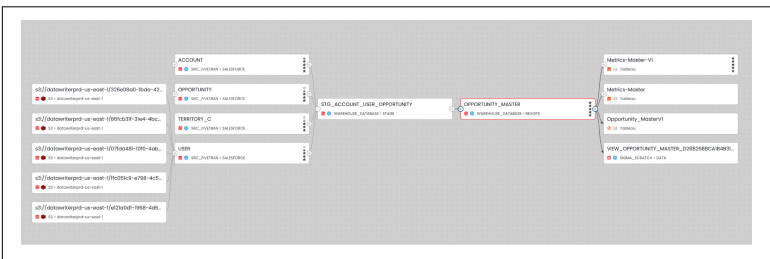


Figure 7-3. Using lineage to understand the flow of your data from source to consumption

With this lineage, it's easy to perform impact analysis to avoid breaking dashboards or reports, because you can see what depends on a dataset before you change it. You can also perform root cause analysis, tracing confusing or incorrect values back to their source. And for governance, lineage provides visibility into how sensitive data flows through the system, supporting compliance efforts and data quality reviews.

Quality and Governance Metadata

Coalesce Catalog brings in quality and governance signals so you can make informed decisions about what data to trust and how to use it. If your environment includes testing or monitoring tools, Coalesce Catalog can surface quality results directly alongside your assets. These might include test outcomes, data freshness indicators, or known issues, and they can immediately give users confidence in (or a healthy skepticism about) what they're looking at.

Stewards or admins can certify trusted datasets by selecting the green checkmark in the detail pane of any table or view; such a visible badge allows users to easily identify trustworthy data and gain confidence in the data they are using. At the same time, deprecated assets can be clearly flagged, allowing users to reach out for better guidance or find a certified table. These labels steer teams toward the sources of truth and can be applied to any asset, as shown in [Figure 7-4](#).

You can also flag sensitive data, like personally identifiable information (PII), either automatically or manually ([Figure 7-4](#)). Coalesce Catalog may detect fields like email or ID numbers based on pattern recognition, and then confirmation by data stewards will ensure accuracy. These PII flags appear in search, filters, and lineage, helping teams protect and handle data responsibly.

This is another case where lineage can support governance workflows. Since teams can track where sensitive data originates from and where it flows, it is easier for them to manage policies like data retention or privacy compliance. Seeing the full path of a dataset builds trust and enforces transparency.

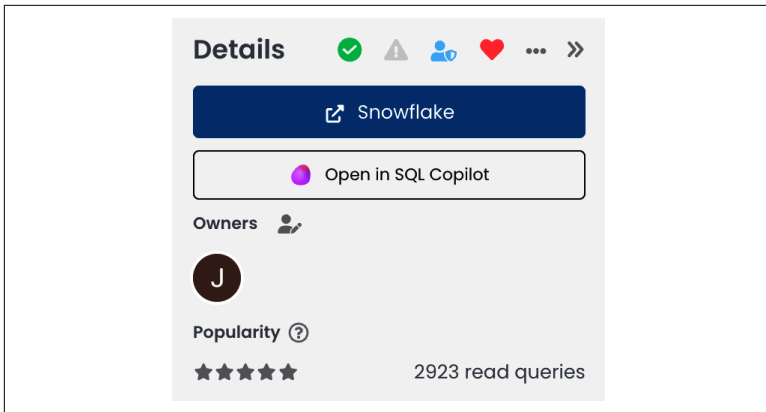


Figure 7-4. Labeling a table as certified (green checkmark), containing PII (blue person), or a personal favorite (red heart)

By unifying documentation, lineage, quality signals, and governance context, Coalesce Catalog becomes a single source of truth for how your data works—and how it should be used. This in turn promotes self-service discovery and can be a core driver in collaboration for all teams working with your data.

Collaboration and Data Discovery

Fostering a data-driven culture takes more than cataloging assets. It requires making data easy to find, understand, and trust. Coalesce Catalog is designed to encourage collaboration and simplify discovery so that everyone, from engineers to analysts to business users, can participate in growing the organization's knowledge base.

Collaborative Features

Documentation in Coalesce Catalog is treated as a team sport. Anyone can contribute, ask questions, or surface issues, making it easier to crowdsource knowledge and keep your documentation current.

Every data asset includes a comments section where users can ask questions or start discussions. These threaded conversations are visible to others and persist, reducing the number of repeated questions and giving future users valuable context. For example, a business user might flag an unexpected number, and an analyst could reply with an explanation or link to related data.

When documentation is missing, users can request it directly from the asset. This sends a notification to the assigned owner or data team, creating a lightweight workflow to close documentation gaps. Once the request is fulfilled, the user is notified, keeping the loop tight and productive.

If there's a problem with a dataset, like unexpected nulls or incorrect values, users can report an issue directly from the Catalog. This keeps issues tied to the source instead of scattered in Slack threads or buried in emails. When the issue is resolved by your data or engineering team, it can be closed with context for others to review later.

The Slack integration plays a major role in bringing Coalesce Catalog into your team's daily workflow. Notifications can be routed to Slack channels, and users can search the Catalog directly in Slack via the native Coalesce Catalog assistant. For example, asking "Do we have any dashboards on regional sales?" in a data channel may return a few results with direct links of high relevance. This dramatically reduces the number of one-off data questions.

These features work together to create a shared sense of responsibility and community around data. Producers document and maintain, consumers ask and learn, and the Catalog becomes more useful and mature over time. But sometimes users need to dig deeper than they can by just asking a question in Slack. In this case, they can use the Catalog's rich discovery features to answer their questions.

Discovery Features

Beyond collaboration, Coalesce Catalog helps users find what they need, even when they don't know exactly what they're looking for. The search experience is tuned to prioritize the most useful results; popular assets, certified datasets, and frequently accessed resources appear higher in results, helping users zero in on trusted sources. If you're new to the Catalog and search for "Sales," the most reliable table or dashboard will rise to the top.

On the homepage or in search results, Coalesce Catalog can highlight trending assets—datasets or dashboards with a spike in usage. This visibility helps teams stay aligned during reporting cycles or major projects. For each asset, you can see who uses it most frequently. This makes it easy to find someone to reach out to when you have questions. Whether you're new to the company or just

working with a new domain, knowing who uses the data helps you find context fast.

The knowledge section serves as a data wiki for your entire data team. It's the place to define business terms like *active users* or *net revenue* and to link those definitions directly to the datasets or dashboards that implement them. If someone searches a business concept, they might land on a glossary page that explains the metric and links out to relevant data assets. This is incredibly beneficial for providing a singular definition for metrics used throughout the organization.

The Catalog browser extension lets you bring Coalesce Catalog into other tools. When you are working in a database UI or business intelligence (BI) tool, the browser extension will show you all of the relevant information documented in the Catalog within the page you are viewing, as shown in [Figure 7-5](#). This provides in-context help without the need to switch tabs.

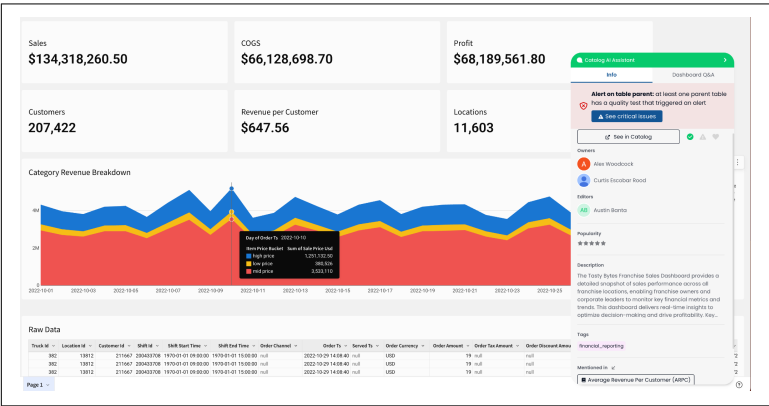


Figure 7-5. The Catalog assistant providing context even when you're not in the Catalog by following you into other tools you already work in

If enabled, the AI assistant adds another layer to data discovery. Users can ask natural language questions like “How do we calculate churn rate?” and get context-rich answers drawn from your documentation. This is especially helpful for nontechnical users who aren't sure what to search for. Over time, Coalesce Catalog learns from these usage patterns and can recommend assets, highlight similar datasets, or identify potential overlaps. This means that every

user's experience is enhanced with the use of AI throughout each touchpoint of Coalesce Catalog.

Using the AI Assistant for Data Discovery and Exploration

The AI assistant in Coalesce Catalog acts like a knowledgeable colleague who is available anytime. Using natural language, anyone can explore the catalog, uncover definitions, trace data lineage, or find the right dashboard without knowing table names or writing SQL.

The AI assistant lets you search the Catalog in plain English, as shown in **Figure 7-6**. Whether you're looking for "monthly revenue by region" or "customer churn rate," it interprets your intent and surfaces the most relevant datasets, dashboards, or glossary terms. It's tuned to your company's language and structure, so it understands business context and not just schema names.

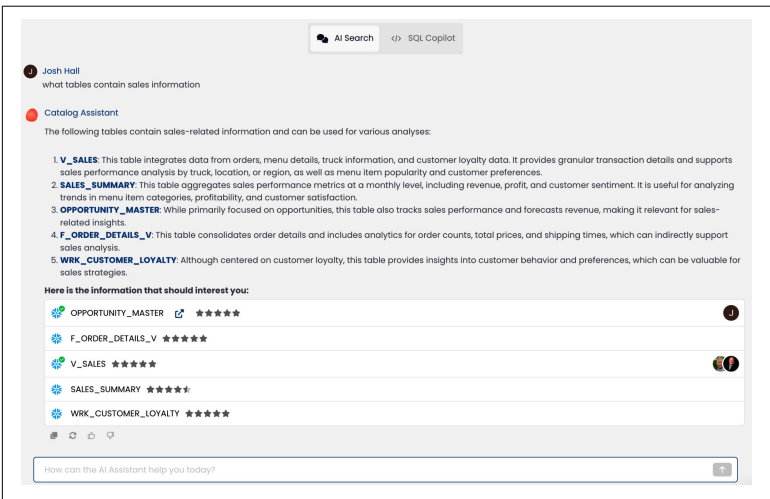


Figure 7-6. Using the AI assistant to search the catalog in plain English

It also acts as an instant Q&A tool. You can ask questions such as "How is churn rate calculated?" and it will respond with the official definition from your knowledge base, along with links to relevant assets. If you're not sure where customer purchase history is stored, you can ask, and the assistant will locate the right datasets for you.

Beyond answering questions, the AI assistant also connects the dots. Ask about a dataset, and it might tell you what dashboards use it, who owns it, or what glossary terms are related. If you're already viewing a table, it adapts to you, letting you ask things like "What upstream source does this come from?" or "Are there issues reported on this asset?"

You can access the assistant from within the Catalog interface, whether browsing a table, searching across assets, or chatting via integrated tools like Slack. This allows users to take advantage of all of the rich context stored in the Catalog regardless of which system they're working in, meeting them where they already are. In doing so, the entire data transformation cycle is brought full circle, giving users insight into even the most complex data transformations, whether or not they know how to write SQL themselves.

From Chaos to Clarity

You've seen it all at this point. You know what it looks like to build a modern, high-functioning data practice using Coalesce. From hands-on SQL transformations to catalog-wide collaboration, this guide has walked you through the full lifecycle of managing, documenting, and delivering trusted data at scale.

At its core, Coalesce is designed to take the repetitive, brittle, and fragmented parts of data work and turn them into something structured, powerful, reusable, and automated. You started by learning how Coalesce transformation gives engineers a clean, visual interface for building logic while still generating code that's precise and production-ready. With every node you added, every column you documented, and every relationship you configured, you built something bigger than a data pipeline—you built a process others can follow.

Then, with Coalesce Catalog, you learned how to capture the entire context around your data: where it comes from, how it's used, who owns it, and what it means. No more scattered documentation. No more Slack messages asking what "ARPU" stands for. No more hoping someone wrote it down in a Google Doc two quarters ago. Catalog ties everything together—your lineage, your logic, your glossary—and puts it at your fingertips. And with the AI assistant layered on top, even nontechnical users can find what they need with a simple question.

Together, these two parts, transformation and the Catalog, create a feedback loop. When your transformations are well documented, they're easier to understand. When they're easier to understand, they're easier to trust. And when people trust the data, they use it. That's how you shift your organization from pipeline firefighting to true data enablement.

This isn't about moving data from point A to point B. It's about creating systems people can understand, build upon, and rely on. It's about letting engineers scale their output without adding complexity. It's about letting analysts explore without gatekeeping. It's about empowering business users to make decisions without waiting in line.

As you move forward, here are a few things to keep in mind:

- Build with clarity. Your future teammates will thank you.
- Don't treat documentation as a separate task. When it's part of the workflow, it actually happens.
- Make the Catalog the first stop for every data question.
- Use lineage not just to debug but to build trust.
- Enable business users to answer their own questions with the AI assistant. This gives you more time to work on high-value tasks.

What you build using Coalesce isn't just a set of tables and models. It's a foundation: one that enables faster development, better governance, and a more confident data culture. So whether you're scaling a team of a hundred or holding down the fort as a team of one, Coalesce gives you the structure, automation, and transparency to do more, with less friction.

Now it's your turn to carry that momentum forward. Keep iterating. Keep sharing knowledge. Keep raising the bar for what good data work looks like in your organization.

Now go build something remarkable.

About the Author

Josh Hall is a data professional with experience spanning data engineering, sales engineering, and technical product marketing. He is passionate about helping others learn, explore, and understand the power and impact of data. Josh has spent much of his career turning complex data concepts into practical insights, through hands-on engineering work, customer-facing solution design, and creating educational content for the broader data community. With over half a decade of consulting experience and several more years building and marketing data products, he shares his hard-earned knowledge so others can accelerate and amplify their own data journey.